



Center for Effective Lawmaking

Measuring Partisanship and Representation in Online Congressional Communications

Michael R. Kistner

Assistant Professor of Political Science
Political Science, University of Houston

Michael Heseltine

Postdoctoral Researcher
Sociology, University of Oxford

Robert Alvarez

Assistant Professor
Political Science, Sam Houston State University

Maya Fitch

PhD Candidate
Political Science, University of Houston

Lucas Lothamer

Assistant Professor
Political Science, University of Texas (Tyler)

Elizabeth N. Simas

Associate Professor
Political Science, Texas A&M University

November 2025

Abstract: Social media and the internet have created new ways for representatives to communicate. How have members of Congress responded to these opportunities? We introduce a multi-platform dataset of congressional communications extending back to the onset of the social media era. Using computational language processing, we classify approximately 4.7 million tweets, 2.4 million Facebook posts, and 184,000 email newsletters authored by members of Congress between 2009-2022 based on intended purpose, and scale the partisanship of each message along a continuous left-right dimension. After validation, we demonstrate how our data can be used to study partisanship and representation in the contemporary Congress. Importantly, our data show congressional rhetoric has become more partisan and negative as social media usage has increased. We identify one potential mechanism contributing to this trend: partisanship and negativity receive inflated levels of positive engagement on social media relative to other forms of messaging like credit claiming or constituency service.

Keywords: Representation, polarization, social media, political communication, Congress

Word Count: 9,746

Center for Effective Lawmaking Working Paper 2025-11

Introduction

Communication is a critical component of representation ([Burke 1774](#); [Cohen 1989](#); [Habermas 1968](#); [Mansbridge 2003](#); [Mill 1861](#); [Pitkin 1967](#)). Effective representation requires a dynamic, deliberative sharing of information between the elected official and those they represent in office. Although communication has always played a central role in representative government, developments in both society and technology affect what information gets shared and how.

In this paper, we introduce new data and measures that shed light on how rhetoric by elected representatives has evolved in the era of social media, email, and other online communication. Our focus is Congress, perhaps the most studied representative assembly. In the 1970s, scholars such as Mayhew ([1974](#)) and Fenno ([1978](#)) put a spotlight on congressional communications, arguing that legislators strategically promote themselves to various constituencies for electoral and other reasons. A key insight was that most messaging by legislators aims to achieve one of a few universal goals, such as promoting a personal brand, claiming credit for legislative accomplishments, taking stances on issues voters care about, or adopting a home style.

Fifty years later, there has been both continuity and change in congressional communication. Members today continue to advertise, claim credit, take positions, and connect with constituents. But the way in which they do so differs considerably from earlier decades. Unlike when Mayhew and Fenno were writing, the advent of online communication and social media mean that legislators can reach a vast audience both inside and outside of their district boundaries nearly instantaneously ([Gainous and Wagner 2013](#); [Russell 2021b](#)). Traditional gatekeepers such as parties and legacy media organizations no longer maintain an oligopoly on communication, enabling politicians to connect with the public directly using chosen messages ([Jungherr and Schroeder 2022](#); [Schroeder 2018](#)). Changes in the technological environment have been accompanied by changes in the political environment. Partisanship and polarization both inside and outside of Congress have grown considerably ([McCarty 2019](#)).

This paper advances the study of representation and communication in two ways. First, we introduce new data on communications by members of Congress spanning multiple platforms over a 14 year period (2009 to 2022), covering almost the entire time period that social media has been widely used.¹ The dataset includes over 4.7 million tweets, 2.4 million Facebook posts, and 184,000 email newsletters. A textual scaling model is used to place each message on a partisan spectrum from far left to moderate to far right based on the language used. Similarly, a set of transformer-based classification models pre-trained on social media are used to classify each message into six different categories based on the intended purpose of the message, classifications that encompass over 72% of all communications in our dataset.

Second, we demonstrate the possibilities our data unlock via several applications. In these applications, we chart the development of congressional rhetoric across the social media period. Congressional rhetoric has become more partisan as well as more negative, trends found across all platforms. Possibly contributing towards this development, we find that partisanship, position-taking, and negativity receive much more engagement on social media than other types of messaging (e.g., credit claiming or constituency service).

The paper proceeds as follows. First, we detail our data collection process and the computational methods used to measure partisanship and classify representational content in congressional communications. After a series of extensive validations of the measures, we proceed to the applications described above, examining temporal trends in rhetoric and exploring differences in online engagement. Throughout, we demonstrate how our new dataset enables more comprehensive investigations of elite political communication than previously possible.

¹Both Facebook and Twitter/X first became available for public usage in 2006. For brevity, throughout this paper we refer to the latter platform as Twitter, as the name change to X occurred after our dataset’s timespan. In the conclusion, we discuss recent developments such as platform changes in congressional social media usage

Congressional Rhetoric and Online Communication

The literature on congressional representation has long recognized that members play an active and strategic role in communicating information for various purposes, particularly electoral success (Fenno 1978; Mayhew 1974; Yiannakis 1982). For instance, Fenno (1978) emphasizes how political representatives shape their public image to align with crucial district elements, crafting a distinct “homestyle” to cultivate the trust of their constituents. Mayhew (1974) highlights three activities in particular that members do to enhance their re-election chances: position-taking (establishing stances on political issues); credit claiming (taking responsibility for work done to pass legislation and secure resources for their districts); and advertising (bolstering name recognition, highlighting appearances at events and mentions in the media). More recent work confirms that members of Congress continue to use similar messaging strategies in modern times (e.g., Ban and Kaslovsky 2025; Grimmer 2013; Hunt and Miler 2025; Russell 2021*b*).

But while the underlying purposes of communication may seem similar, the means of communication have changed dramatically. In the contemporary era, the rise of the internet and social media has transformed communication, allowing representatives to reach a wider audience than ever before. Members of Congress are no longer dependent on a franking privilege to directly reach those they represent, and even lowly rank-and-file legislators now have the ability to broadcast views far beyond the boundaries of their districts. Moreover, messages no longer have to pass through newspaper reporters or television anchors before reaching the public (Gainous and Wagner 2013; Russell 2021*b*). The internet and social media allow representatives to reach constituents in a direct and unmediated way. At the same time, the proliferation of information requires information to break through the noise; with the demise of traditional news sources such as local papers (Canes-Wrone and Kistner 2023; Hayes and Lawless 2015, 2017; Peterson 2021), there is no single front page or editorial section equivalent that guarantees messages will be seen by a large, captive audience (Jungheer and Schroeder 2022; Schroeder 2018). Social media has

also enabled two-way communication between the mass public and political elites, allowing each group to influence the discourse of the other ([Barberá et al. 2019](#); [Warner 2023](#)).

Researchers have begun studying how political elites use these new forms of media, although the field is still nascent. As recently as 2017, scholars were referring to online tools of campaign communication – “smartphones, Facebook, blogs, and the like” as “niche communication(s)” ([Frankel and Hillygus 2017](#)). Still, considerable progress has been made in understanding how these tools are used. We view these works as creating at least two major strands of research.

One of these major strands focuses on explaining what factors shape differential usage of social media and online communication, both in terms of the volume of social media usage as well as which types of messages (e.g., position-taking versus credit claiming) members choose to prioritize ([Albert 2020](#); [Cormack 2016a,b](#); [Evans and Clark 2016](#); [Evans, Cordova, and Sipole 2014](#); [Hemphill, Russell, and Schöpke-Gonzalez 2021](#); [Hunt and Miler 2025](#); [Russell 2018a,b, 2021a,b](#); [Scherpereel, Wohlgemuth, and Lievens 2018](#); [Smith and Russell 2022](#); [Straus et al. 2016](#); [Tillery 2021](#)). This work typically considers the political, institutional, demographic, sociological, and other variables that lead politicians to communicate in different ways.

A second major strand of research focuses on how rhetoric has evolved in response not just to the rise of social media, but also to the growing political polarization in American society. While many have studied the polarizing effects of social media usage on the mass public (for a review, see [Tucker et al. 2018](#)), others have focused more specifically on how political elites use social media in polarizing ways, via the language they use and how they discuss political issues online. Research on congressional rhetoric has considered the extent to which lawmakers deploy polarizing or extreme rhetoric ([Ballard et al. 2023](#); [Cowburn and Sältzer 2024](#); [Heseltine 2023, 2024](#); [Kaslovsky and Kistner 2025](#)), negativity ([Macdonald, Russell, and Hua 2023](#); [Macdonald, Hua, and Russell 2024](#); [Russell 2018a](#); [Yu, Wojcieszak, and Casas 2024](#)), and uncivil language ([Ballard et al. 2022](#)). Once again, a common theme in this work is the importance of electoral incentives in shaping member behaviors. Both social media users and donors (categories with some overlap)

reward this type of language with engagement (Ballard et al. 2023; Macdonald, Russell, and Hua 2023) and dollars (Hilden and Kistner 2025; Fu and Howell 2020; Yu, Wojcieszak, and Casas 2024).

While this research has improved our understanding of how representatives communicate, most of this work has been hamstrung by three key limitations. The first limitation, common to almost all of the above-cited research, is a focus on short time periods. Due in part to the difficulties in collecting and cleaning communication data, most research uses at most a few years of data, which can lead to inferential issues.² For instance, studies using data from a single session (Hemphill, Russell, and Schöpke-Gonzalez 2021; Yu, Wojcieszak, and Casas 2024) have found communication differences between Democrats and Republicans, which are attributed to one party being in the majority and the other in the minority. But Democrats and Republicans differ from each other in many ways, making it impossible to separate partisan differences from majority-minority differences when analyzing just a single session’s worth of data.

Besides avoiding problems such as these, longer time spans are desirable for another reason. It’s unclear whether congressional communication has stayed largely constant or evolved as social media usage and technology have changed. Particularly given the speed of developments in online communication, a major concern is the temporal validity of findings in this area (Munger 2023). Assessing how stable conclusions are requires data covering longer periods of time.

A second issue with most existing research is examining communication on a single platform.³ The audiences members speak to when posting on a social media platform like Twitter/X – where messages are seen by a heterogeneous mix of political enthusiasts, journalists, interest group members, fellow politicians, and more – look very different from the recipients of email newsletters, which are targeted more directly towards constituents.⁴ Recently published research

²Section A of the Supplemental Materials displays a table of published political science articles over the past twelve years that analyze text from online Congressional communication data (Twitter/X, Facebook, or e-newsletters). Out of these 30 articles, 9 used one year of data, 8 used two years, 4 used three years, 3 used four to six years, and 6 used more than six years of data.

³Again referencing Section A of the Supplemental Materials, only 4 of 30 published articles studying online congressional communication examined more than one platform.

⁴On that latter point, offices sometimes require individuals sign up for e-newsletters using a zip code, to confirm

demonstrates that these different audiences matter. Members vary in terms of how much they post on Facebook versus Twitter/X (Blum, Cormack, and Shoub 2023). Furthermore, the partisanship of member speech can vary by venue, as sometimes members appear more partisan when measured using communication in one form versus another (Green et al. 2024). Other hypotheses researchers are interested in testing may be platform-specific, and demonstrating similarities or differences across platforms can provide deeper insight.

The third issue is that the current system – where research teams individually download, clean, and prepare different versions of similar datasets – is inherently wasteful and slows the pace of scientific progress. Having a central repository of easily accessible, ready-to-use data allows researchers to spend their time developing and testing theories of political communication, not repeating time-consuming data work that has been done many times over. To build on this point, having a single commonly-used dataset ensures that similar cleaning and sample inclusion decisions have been made. Idiosyncratic data processing decisions (that may not be immediately obvious to readers) can be eliminated as possible explanations when differing results emerge, making comparison of results more transparent.

For these reasons, the study of communication and representation stands to benefit enormously from a single publicly available dataset with multimodal and longitudinal data, classified by representational purpose and scaled according to the partisan positioning expressed via the message. In the remainder of the paper, we describe our efforts to accomplish exactly that.

Measuring Partisanship and Representation

To address these limitations and advance future research, we create the *Scaled and Classified Congressional Communication (SCCC) dataset*, a multimodal dataset spanning the years 2009 to 2022. The dataset includes posts on the two most-used social media platforms by members of

constituency residency.

Congress (Twitter/X and Facebook) as well as email newsletters, a common form of communication members used primarily to reach constituents. In addition to texts of communications, we possess auxiliary variables such as social media engagement metrics (likes, retweets, shares, etc.).

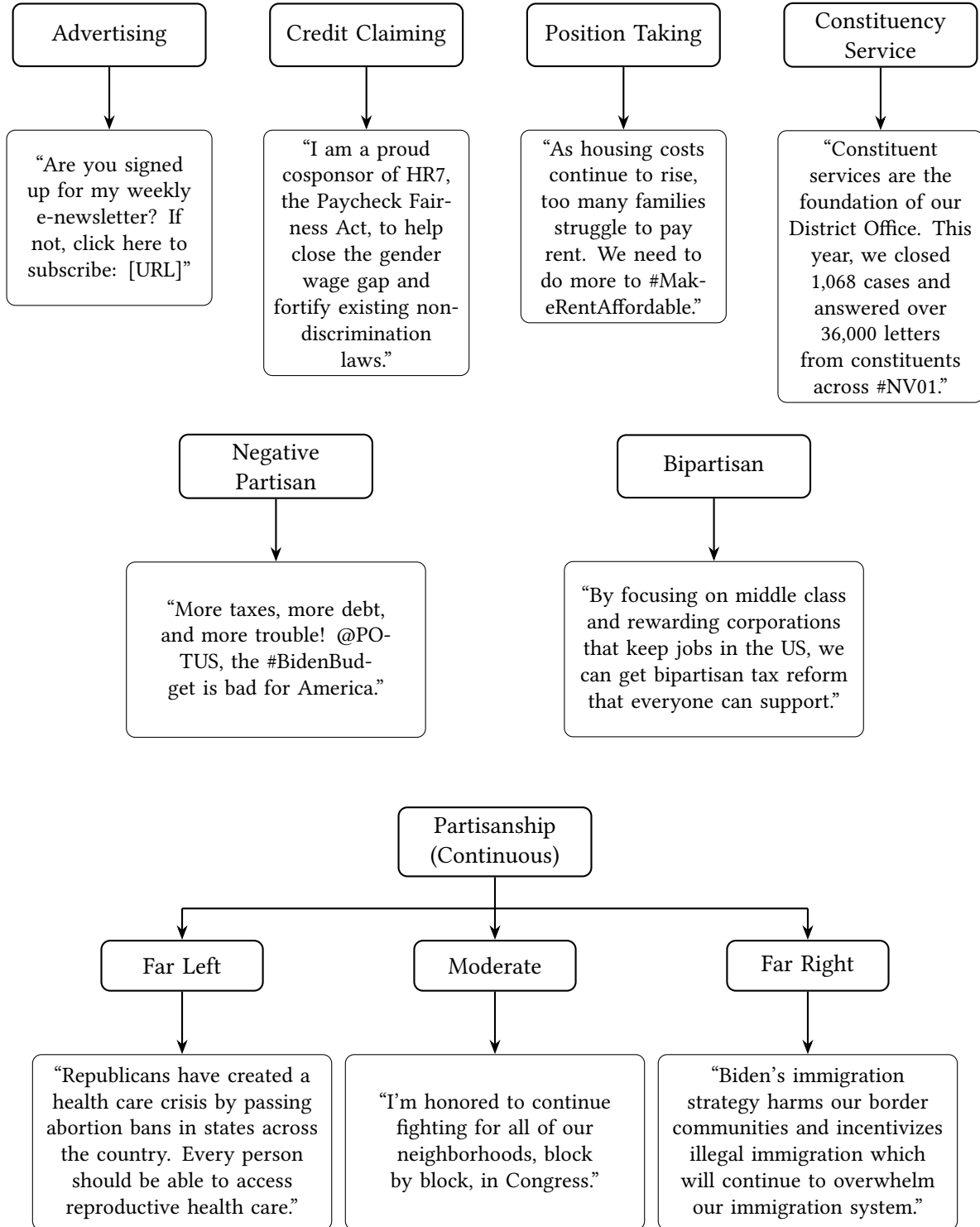
Measurement schemata are shown in Figure 1, along with corresponding example messages. The representational categories come directly from Mayhew (1974) and Fenno (1978). These categories are *Advertising* (“any effort to disseminate one’s name among constituents in such a fashion as to create a favorable image”), *Credit Claiming* (“generat[ing] a belief...that one is personally responsible for causing the government, or some unit thereof, to do something that the actor...considers desirable”), *Position Taking* (“a judgmental statement on anything likely to be of interest to political actors”), and *Constituency Service*, which corresponds to two components of home style: the “presentation of self” in the district (meetings with constituents, attendance at social events, etc.), and “allocation of resources” to the district (helping constituents with obtaining passports, social security checks, or internships in public office).⁵

The partisanship categories mirror those studied by Russell (2018b); specifically, we identify both *Negative Partisanship* (an attack on the policies and politicians of the opposing party) and *Bipartisanship* (advocating the value of bipartisan collaboration) in messages.⁶ In addition, we estimate the *Partisan Orientation* of each message on a scale that ranges from -1 (most Democrat-leaning) to 0 (nonpartisan) to 1 (most Republican-leaning).

⁵The third component of home style, “Explaining Washington Activity”, fits closely in our Position Taking category.

⁶Russell (2018b) also studies a third partisanship category, *Positive Partisanship*, messages that “signal favoritism or support for one’s own party [or one’s] party’s candidates” (p. 703). We omit this category as there were few messages classified as positive partisan by our human coders and (consequently) the accuracy of the algorithmic classifications was much lower for this category than the others.

FIGURE 1: CLASSIFICATION AND SCALING SCHEMATA WITH EXAMPLE TEXT



The categories are neither mutually exclusive nor collectively exhaustive. A single tweet could, for instance, advertise (“As the representative for TX-23, ”...), position-take (“securing the border is one of my top priorities.”), credit claim (“This is why I’m sponsoring legislation...”), and make a negative partisan attack (“The Biden border crisis must be stopped!”). It can also be none of the above, wishing (for example) followers a Merry Christmas or Happy Holidays.

While we choose these categories in part based on congruence with existing research, these categories also each represent critical components of the information voters require in order to hold legislators accountable. A useful theoretical framework to think about representation is the *principal-agent model of political accountability* (Ashworth 2012; Besley 2006), in which voters select politicians with aligned preferences whom they believe will do the best jobs of advancing their interests. As a consequence, politicians motivated by re-election take actions that correspond to voter preferences and benefit constituents, in order to maximize their chances of remaining in office. This relationship requires that voters know who their representatives are (which *advertising* facilitates), how they vote for or against specific policies (*position-taking*), and the actions they take to aid voters (*credit-claiming, constituent service*). *Bipartisan* and *negative partisan* messages send signals to voters about whether representatives will work with or oppose members of the other party, also critical pieces of information in a polarized age.

Data Collection

The communications data come from multiple sources. For tweets and Facebook posts, data were downloaded directly on a rolling basis beginning in 2018.⁷ Tweets were first downloaded using the V2 API endpoint and then later using the Academic Twitter API. For Facebook, posts

⁷Our Twitter/X and Facebook data includes official, campaign, and personal accounts publicly associated with members of Congress. All account names/handles were hand-collected by one of the authors, with periodic updates to include new Members and check for newly created accounts. Content from posts and accounts deleted before the date of collection were, by definition, not retrievable. In the case of Facebook accounts, not all accounts were created as official accounts. As Crowdtangle only facilitated the collection of data from specifically-created flagged account types and not personal pages, retrieval of data from some accounts was not possible, an issue which mostly impacted accounts from earlier time periods.

were downloaded using the CrowdTangle Platform. For newsletters, data come from the publicly available www.DCInbox.com repository of email newsletters collected and cleaned by Cormack (2017). For newsletters, the data are available dating back to 2010; in the case of Twitter/X and Facebook, the data are available dating back to 2009. The dataset currently spans through 2022, although we aim to make periodic future updates of the dataset to broaden the timespan and enable study of contemporary congressional communication.

Table 1 displays the total number of tweets, Facebook posts, and email newsletters included in the dataset, listed by biennial legislative session. In total, the data consist of 7,827,972 unique communications from 1,034 US senators and representatives. As the table shows, the volume of online communication has grown, although some of the social media differences in the pre-2018 years may be due to the deletion of accounts. For Twitter and Facebook, the table thus displays lower bounds on the total number of users as well as tweets and posts from this time period. On the other hand, the number of tweets and posts from the median member has gone up considerably over this time period (over a fourfold increase for each social media platform), suggesting that the increase in social media usage over this time period is not merely an artifact of missing accounts.

Scaling Procedure

To scale speech as more or less partisan, we adopt the same general strategy as Gentzkow, Shapiro, and Taddy (2019), who measure partisanship in Congressional floor speeches. Following their lead, we define speech as partisan on the basis of how strongly it identifies the party of the speaker. Extreme partisan speech is used almost exclusively by members of one or the other party, while moderate speech is used by both. For instance, an individual who uses terms such as “border crisis” when discussing immigration is likely to be a Republican, while the utterance of “pathway to citizenship” means the speaker is likely to be a Democrat.

We opt to scale speech *partisanship* rather than speech *ideology* because we view partisanship

TABLE 1: ONLINE COMMUNICATION BY MEMBERS OF CONGRESS (2009 - 2022)

Medium	Session	Years	Percent Using	# (Median Member)	# (Total)
Tweets	111	2009-2010	–	–	71,061
	112	2011-2012	82.0	444.0	342,597
	113	2013-2014	89.4	811.5	578,740
	114	2015-2016	91.1	901.0	669,638
	115	2017-2018	93.8	1,153.5	865,703
	116	2019-2020	96.7	1,574.5	1,089,424
	117	2021-2022	97.8	1,676.0	1,100,253
Facebook Posts	111	2009-2010	–	–	42,669
	112	2011-2012	69.3	187.5	147,617
	113	2013-2014	80.8	261.5	192,816
	114	2015-2016	86.8	512.0	324,306
	115	2017-2018	90.7	722.0	447,973
	116	2019-2020	95.0	891.0	587,363
	117	2021-2022	96.9	1061.5	677,376
Newsletters	111	2009-2010	–	–	8,076
	112	2011-2012	88.8	26.0	21,620
	113	2013-2014	90.3	27.0	22,355
	114	2015-2016	92.2	30.0	23,749
	115	2017-2018	91.9	32.0	24,703
	116	2019-2020	92.4	43.0	30,406
	117	2021-2022	88.5	39.0	29,383

Note: The table displays the usage of tweets, Facebook posts, and email newsletters by members of Congress in our dataset. Percent Using and Number by Median Member are excluded for the 111th session, as the data do not span the full session.

as a broader concept inclusive of more relevant features when evaluating rhetoric than ideology. Similar to Converse (1964), we conceptualize ideology as a set of co-occurring (i.e., constrained) issue beliefs. We conceptualize partisanship similar to how the Michigan School authors describe partisanship in chapter 10 of *The American Voter* (Campbell et al. 1960), which encompasses issue views but also friendliness or hostility to various political figures or socio-demographic groups. Partisan speech may include not just the discussion of policy issues, as in the examples in the prior paragraph, but also non-policy topics, such as individual politicians (e.g., “crooked Hillary”). In an era where certain members of Congress have chosen not to hire any policy staff so they can employ more communications staffers (Vesoulis 2021), it is necessary to measure members’ partisan positioning in a way that incorporates not just issue stances but also broader partisan language. Partisanship can also be inferred from the way in which politicians talk. In an era of increasing educational polarization (Grossmann and Hopkins 2024), the degree of formality and refinement may provide important signals about how a politician presents themselves.

To estimate the partisanship of member language, a class affinity scaling model (Perry and Benoit 2017) is fit using the words contained in the tweets, Facebook posts, and newsletters members share. Though the party identification of the speaker is used in training the algorithm, the procedure is better described as scaling than supervised classification since it models speakers having a continuous affinity towards classes (e.g., parties) rather than simply predicting membership in a binary class.⁸ The model is estimated separately for each session of Congress, to account for changes in the partisan valence of language over time.

In the model, affinity towards party (partisanship) is parameterized as $\pi_r \in [0, 1]$, the proba-

⁸We chose the class affinity scaling model based on theoretical motivation (as a continuous measure of partisanship rather than ideology), interpretability, and computational efficiency.

In Supplemental Section B, we discuss the choice of a scaling algorithm and compare the results to alternative scaling options such as Wordscores (Laver, Benoit, and Garry 2003), Wordfish (Slapin and Proksch 2008), Correspondence Analysis (Greenacre 2007), or using the predicted probabilities of a supervised classifier such as Naive Bayes. As shown in Figure B.1, the class affinity scaling model produces member-level estimates that correlate more strongly with a composite of different positioning measures (based on roll call votes, contributions, presidential support, and website text) than the other algorithmic options.

bility that for any given token of speech W_i the underlying orientation of the speaker is $U_i = r$, or Republican. The probability that the speaker’s underlying orientation is Democratic ($U_i = d$) for a given token of speech is thus $1 - \pi_r$. The orientation of a speaker for each token $i = 1, \dots, n$ determines the probability of a specific word w being used:

$$\Pr(W_i = w) = \pi_r \Pr(W_i = w | U_i = r) + (1 - \pi_r) \Pr(W_i = w | U_i = d).$$

The partisanship for any given tweet, Facebook post, or newsletter is the expected proportion of time the underlying orientation is r versus d , which can be calculated as $\pi_r = \mathbb{E}\left\{\frac{1}{n} \sum_{i=1}^n U_i = r\right\}$. For interpretability sake, we rescale the resulting probability so it ranges from -1 (most Democrat-leaning) to 0 (neutral partisanship) to 1 (most Republican-leaning), calculated as $2\pi_r - 1$. This rescaling we refer to as the *Text Partisanship Score*. This measure can be folded, i.e., $|2\pi_r - 1| \in [0, 1]$, so that higher values indicate more extreme partisan rhetoric regardless of the speaker’s party. This measure we refer to as the *Text Partisan Extremity Score*.

We validate our Text Partisanship Scores in several ways. First, we aggregate the text-level scores by taking the average for each member of Congress in our dataset, and compare our scores to four other common measures of positioning by members of Congress: DW-NOMINATE (Poole and Rosenthal 2000) based on members’ roll call votes, CF scores (Bonica 2014) based on members’ fundraising, platform positions (Meisels 2025) based on issue statements on members’ websites, and presidential support scores (Edwards 1985), based on how frequently members vote for or against the stated preferences of a co-partisan or opposite-party president.⁹

Figure 2 compares the within-party correlations of our aggregated Text Partisanship Scores with these four alternatives, as well as a composite *positioning dimension* constructed by applying

⁹Three of these measures (DW-NOMINATE, CF scores, and website platform positions) are best described as measures of ideology rather than partisanship. Thus we ex ante expect some differences between these measures and our measure of partisanship. That said, in recent years in American politics there is a tight connection between ideology and partisanship among elites, so we expect (and find) that these ideology measures will have a considerable correlation with our text-based partisanship measure.

Principal Components Analysis to all of the measures and extracting the dimension which explains the most variance.¹⁰ Each of these measures captures member positioning along a different domain, and though there should be shared patterns there will also be domain-specific differences. As Tausanovitch and Warshaw (2017) describe in their comparison of different ideology estimation procedures, “existing measures do not measure the same underlying dimension...The performance of these measures varies across parties, with no measure clearly dominant” (p. 183).¹¹

As Figure 2 demonstrates, our Text Partisan Scores are strongly correlated with the shared positioning dimension and substantially correlated with the individual positioning measures.¹² The first row of each plot shows the Pearson correlations between our member-level Text Partisan Scores and each of the other measures. The first entry is the correlation between our scores and the common dimension uncovered via PCA. For both Democrats and Republicans, the correlation is moderately high both in absolute terms and relative to the other positioning measures. For Democrats, only the Presidential Support score has a slightly higher correlation (by 0.03). For Republicans, DW-NOMINATE and CF Scores correlate modestly higher with the shared dimension (by 0.17 and 0.06 respectively), but the Text Partisan Score has a higher correlation with the shared dimension than either the website platform scores or the presidential support scores.

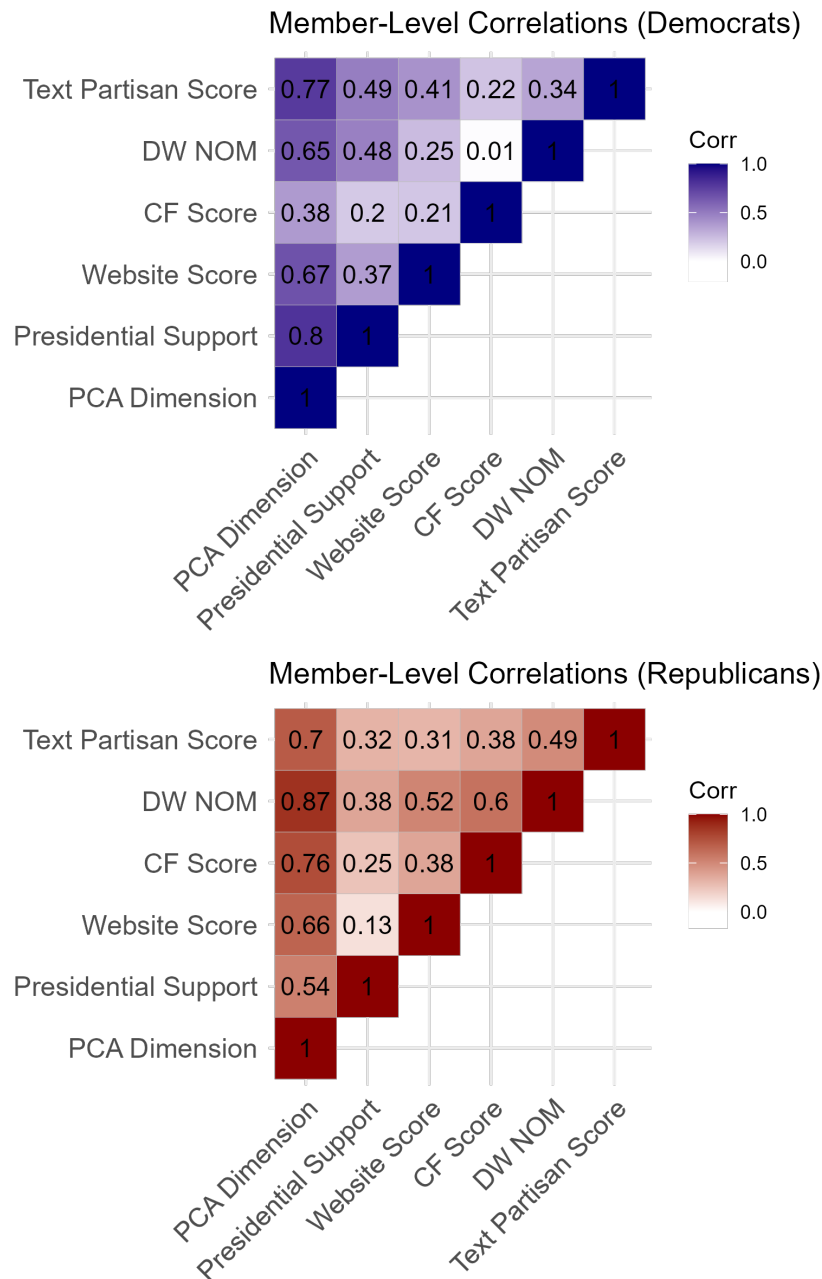
Our Text Partisan Scores also correlate well with the individual measures themselves, ranging from 0.22 (with CF scores for Democrats) to 0.49 (DW-NOMINATE scores for Republicans).

¹⁰The cross-party correlations are much stronger than the within-party correlations shown in Figure 2. For example, our Text Partisan Scores have an overall correlation (i.e., across members of both parties) with DW-NOMINATE scores of 0.91 and an overall correlation with CF Scores of 0.92.

¹¹Tausanovitch and Warshaw frame much of their study in terms of which ideology measures best predict members’ roll call voting. Our goal is to characterize rhetoric, not members’ roll call voting. That said, a possible extension would be to use machine learning methods to infer roll call voting based on member rhetoric, similar to how Bonica (2018) extends his CF score measurements by using machine learning to predict DW-NOMINATE scores based on the contributions members receive.

¹²In many cases we believe our scalings have more facial validity than other measures. For example, members of “The Squad” such as Alexandria Ocasio-Cortez, Ilhan Omar, Ayanna Pressley, Rashida Tlaib, Jamaal Bowman, and Cori Bush are popularly described as some of the most liberal Democrats in Congress. Out of 447 House Democrats in our dataset, these members rank between the 15th and 170th most left-leaning. By first dimension DW-NOMINATE scores, these members rank between the 282 to 379th most liberal House Democrats, less liberal than the average Democrat.

FIGURE 2: COMPARING TEXT PARTISANSHIP SCORES TO EXISTING MEMBER POSITIONING MEASURES



Note: The figure displays the within-party member-level correlations between our partisanship measure (Text Partisan Score) and other commonly-used measures of member ideology and partisanship. Numbers and shading of each cell shows the Pearson correlation.

Notably, some of the non-Text Partisanship measures correlate much more weakly with the other individual measures, including a near-zero correlation between CF scores and DW-NOMINATE scores among Congressional Democrats (for more on this, see [Barber 2022](#)).

The online supplemental materials contain numerous other forms of validation. First, we consider how strongly our partisanship scores correlate with each other if we estimate member partisanship separately using each of the three communication types. As Figure [C.1](#) shows, these within-party correlations are quite high, ranging from 0.51 to 0.82. Second, using a random sample of 3,000 anonymized and human-coded messages, we compare the scalings produced by the class affinity model to human perceptions of the partisan orientation of the member. As Figure [D.1](#) shows, there is almost a linear relationship between these human perceptions and the scaling output. Third, we consider whether our scalings are reflective of observable changes in politicians’ orientations. In Supplemental Section [E](#), we examine changes in our Text Partisanship Scores among members of Congress widely perceived to have evolving partisan loyalties during the time period of our scores. As the section documents, each member’s by-session Text Partisanship Score largely corresponds to public discussion of their changing partisan inclinations. Finally, to establish facial validity and provide researchers intuition as to what types of messages our scaling procedure associates with different positions, Figure [F.1](#) displays the top 20 words most strongly associated with being classified as Democratic rhetoric (a Text Partisanship Score of -1 to -0.5), Republican rhetoric (score of 0.5 to 1), or nonpartisan rhetoric (score of -0.5 to 0.5).

Classification Procedure

To classify messages into the six binary categories displayed in Figure [1](#), trained PhD-level researchers read and coded approximately 43,000 messages for binary membership (yes or no) into each of the categories. These messages were then used to train a series of transformer-based classification models (one for each category) using the BERTweet language model ([Nguyen](#),

Vu, and Tuan Nguyen 2020), a RoBERTa variant pre-trained on English-language tweets.¹³ A comprehensive description of our manual annotation and algorithmic classification procedure is provided in Section G of the Supplemental Materials.

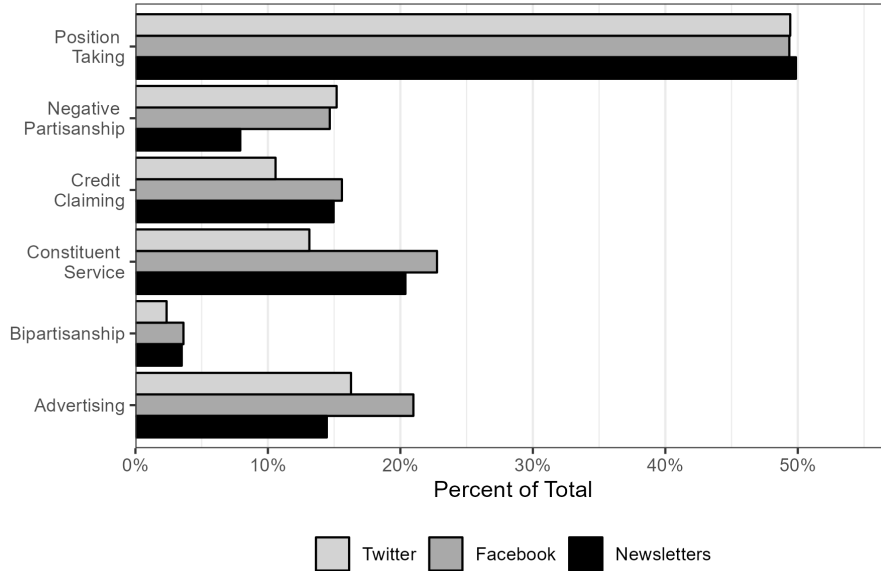
Out-of-sample classification performance metrics for all platforms (accuracy, balanced accuracy, and macro F1) are displayed in Table G.1 in the Supplemental Materials, while more detailed classification metrics are provided in Section H of the Supplemental Materials. The macro F1 for each category of tweets ranged from 0.83 to 0.95, values indicating strong performance. Classifier performance on the Facebook data was substantively identical, with macro F1 scores ranging from 0.81 to 0.92. The models showed marginally weaker performance on the newsletter data, reflecting the more freeform nature of the text compared to social media, but still well within acceptable levels (macro F1 between 0.77 and 0.93).

Figure 3 shows the percentage of all tweets, Facebook posts, and newsletter sentences that can be classified in each category. Across all data sources, position taking is by far the most common type of communication that we observe. This combined with the prevalence of advertising and credit claiming show that while congressional communication has changed, the basic framework of Fenno and Mayhew is still relevant and applicable today. Overall, 71% of all tweets, 76% of all Facebook posts, and 72% of all newsletter sentences can be classified into at least one of our six categories.

Again, we validate the resulting classifications in multiple ways. First, we consider whether members’ messaging on social media and in email newsletters – as measured using our classifications – is related to their behavior in office. We collect data on three outcomes – the number

¹³An alternative to the supervised learning procedure described here would be to use a large language model to perform classifications, either using one-shot learning or fine-tuning a model. Such an approach is less than optimal here, for multiple reasons. First, using a proprietary LLM for such a task would be prohibitively expensive given the approximately 10 million messages that would need to be classified. Second, the classifications themselves would be subject to change based on updates to the LLM, which occur regularly, producing instability and limiting replicability. Despite this, we assessed the relative classification success on a small subset (2,000 tweets) of our data using a fine-tuned GPT-4o mini model. Differences in classification were minimal. As also shown by Heseltine and von Hohenberg (2024), the difference in downstream results based on LLM or transformer-based classification is negligible, at a fraction of the expense or computational resource requirements.

FIGURE 3: FREQUENCY OF MESSAGE TYPES, BY PLATFORM



Note: The figure uses bars to displays the percent of all tweets, Facebook posts, and newsletter sentences classified into each of the six main categories displayed in Figure 1 by the BERTweet models.

of bills introduced by the member, the percent of bills cosponsored by the member that have an opposite-party sponsor, and the number of times they mention a local area in their district (a town, neighborhood, location, etc.) during floor speeches. The bill (co)sponsorship variables are downloaded directly from www.govtrack.us/congress/members/report-cards/, while the local mentions during floor speeches variable comes from Ban and Kaslovsky (2025).

To evaluate whether these in-office behaviors predict messaging, we estimate a series of regression models with the behavioral measures as independent variables, and the percent of a member’s messages (aggregated across platforms) as the dependent variable. We expect members with more bill introductions to do more credit claiming; members with more cross-party cosponsorships to discuss bipartisanship more and make fewer negative partisan attacks; and members who mention localities in floor speeches to reference constituent service more often.

The results of these regression models are displayed in Table 2. For each of the expectations we test, we estimate two versions of the model, one with no control variables and one controlling

TABLE 2: ONLINE MESSAGING AND BEHAVIORS IN CONGRESS

	Credit Claiming		DV: Percent of Messages Classified as...				Constituent Service	
	(1)	(2)	Bipartisanship (3)	(4)	Negative Partisanship (5)	(6)	(7)	(8)
Bills Introduced	0.060*** (0.018)	0.068*** (0.017)						
Percent Opposite-Party Cosponsored Bills			0.108*** (0.011)	0.080*** (0.011)	-0.360*** (0.023)	-0.311*** (0.026)		
Local Mentions in Floor Speeches							1.088** (0.377)	0.777* (0.363)
Num. Obs.	2,099	2,054	2,099	2,054	2,099	2,054	1,503	1,479
Session	113-117	113-117	113-117	113-117	113-117	113-117	111-114	111-114
Controls	N	Y	N	Y	N	Y	N	Y
Party-Session FEs	Y	Y	Y	Y	Y	Y	Y	Y

Note: Estimates are from OLS regression models. Standard errors are clustered by member. * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

for a variety of member characteristics as well as electoral and institutional contexts.¹⁴ As the table shows, all expectations are borne out in the data. At least in aggregate, what members say is what members do.

As with the scaling, we also establish the facial validity of our classifications by creating keyword plots, showing which words are most strongly associated with classification into each of the six categories. Figure F.2 in the Supplemental Materials shows the top 20 words for each category. The words in each category accord with expectations. For example, constituency service messages are associated with words that imply district visits ("town", "hall", "meeting"), aiding constituents ("assistance", "constituents"), and securing money for local projects ("funding", "grant"). Similarly, messages that attack the other party feature the names of top leaders of each party ("trump", "biden", "pelosi") or divisive issues ("obamacare", "border").

¹⁴All control variables originally from the Center for Effective Lawmaking's public data. Full set of estimates including for control variables shown in Table I.1.

The Evolution of Online Congressional Communication

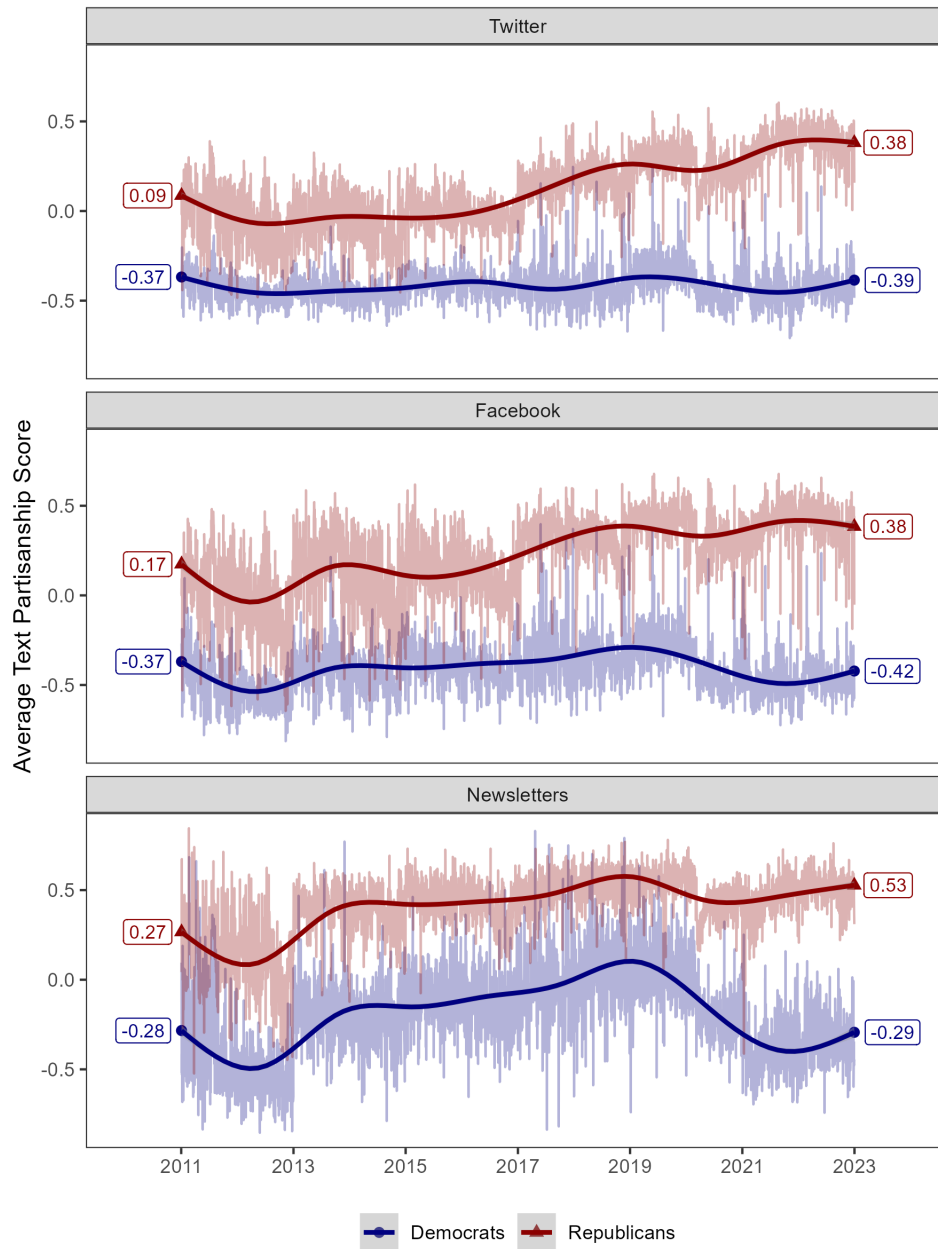
With our new data, the first question we consider is how congressional rhetoric evolved in the age of social media. Has messaging on social media and in newsletters become more polarized? In other domains, such as roll call voting, the congressional parties have been polarizing for decades (Poole and Rosenthal 2000). In our communication-focused context, a growth in polarization would mean that Democrats and Republicans speak in increasingly different ways. There would be a divergence in the issues, people, and groups they discuss as well the words they use to describe them, focusing on partisan topics more and nonpartisan topics (e.g., helping constituents, advertising appearances back home) less.

To evaluate trends in partisan rhetoric across time, Figure 4 plots the average Partisan Extremity Score by calendar day for all three forms of communication and each of the two major political parties across all sessions of Congress for which we have full data (the 112th to 117th Congresses). A smoothed GAM regression line is fit to each of the time series, to flexibly capture changes across the period.

As can be seen in Figure 4, the partisan divide in Congressional rhetoric has grown steadily over this time period across all three forms of communication studied. In 2011, the difference between the language used by Republicans and Democrats on social media was modest, ranging from 0.46 points on the 2-point scale (on Twitter) to 0.55 points (newsletters). By the end of 2022, the difference had almost doubled, ranging from 0.77 points (on Twitter) to 0.82 points (newsletters).

One noteworthy feature of this polarization in congressional rhetoric is that the growth in partisanship is disproportionately driven by Republicans. On the two social media platforms, the partisanship of Democratic speech is mostly flat across this time period, increasing only very slightly. In newsletters, there's a brief period of depolarization from the 113th session of Congress to the 116th session of Congress.

FIGURE 4: PARTISAN SPEECH IN CONGRESSIONAL SOCIAL MEDIA AND NEWSLETTERS (2011 - 2022)



Note: The figure displays the daily average Text Partisanship Score for tweets, Facebook posts, and newsletter sentence bigrams (higher values indicate more Republican rhetoric). Dark lines indicate smoothed GAM regression lines of best fit. Trends shown separately for Democrats (blue; endpoints denoted with circles) and Republicans (red; endpoints denoted with triangles). Numbers denote the average Text Partisanship Score at the beginning and end of each trendline.

How do these trends in partisan rhetoric compare to other forms of polarization? To answer this question, we estimate a series of regression models comparing trends in rhetorical polarization to polarization measured using roll call votes and campaign contributions. For rhetorical polarization, we use the member-session level *Text Partisan Extremity Score* described above, aggregated across all three communication platforms. For roll call polarization, we use (folded) Nokken-Poole scores, a dynamic version of the DW-NOMINATE scores that allows for more sudden change in member ideology than DW-NOMINATE scores themselves do (Nokken and Poole 2004). For campaign contribution polarization, we use (folded) dynamic CF scores (Bonica 2014).

Using these three extremity variables, we estimate within-party OLS regression models of the following form:

$$\frac{\text{Extremity}_{it}}{\sigma_{t=0}} = \alpha + \beta \text{Session}_t + \mathbf{X}_{it}\gamma + \epsilon_{it},$$

where i indexes members and t sessions. To make the results more directly interpretable, Session is recentered so the first session in our analysis (the 112th) is coded as 0 (and the 113th as 1, the 114th as 2, and so on), and the dependent variable is rescaled by dividing by the standard deviation of the extremity measure in the first session. Doing so means the coefficient (β) on the recentered session variable represents the average per-session change in extremity as measured in units of the within-party SD during the first session. As the independent variable is measured at the session-level, we control for two party-session confounding variables: whether the member's party is in the majority in their chamber, and whether there is a co-partisan president. Standard errors are clustered at the member level.

The results are shown in Table 3. The first two columns show the estimates for our Text Partisan Extremity Measure for Democrats and Republicans respectively. After aggregation and including the control variables, there is no statistically significant trend in partisan rhetoric for Democrats. On the other hand, there is a large and significant increase for Republicans. The esti-

TABLE 3: COMPARING POLARIZATION IN RHETORIC TO ALTERNATIVE FORMS

	DV: Partisan/Ideological Extremity					
	Rhetoric		Roll Call Votes		Contributions	
Session	-0.041 (0.022)	0.417*** (0.018)	0.052** (0.019)	0.046** (0.017)	0.165*** (0.018)	0.081*** (0.018)
Majority Party	0.232** (0.076)	0.270*** (0.064)	-0.216*** (0.063)	-0.006 (0.059)	-0.095 (0.060)	-0.042 (0.059)
Copartisan President	0.666*** (0.077)	0.166*** (0.049)	0.074 (0.042)	0.001 (0.036)	0.059* (0.028)	-0.219*** (0.036)
Num.Obs.	1,487	1,630	1,487	1,630	1,487	1,630
Party	Democrats	Republicans	Democrats	Republicans	Democrats	Republicans

Note: Estimates are from OLS regression models. Standard errors are clustered by member. * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

mate suggests that partisanship increased by approximately 0.42 standard deviations per session, implying that by the final session in our data (2021-2022), the average Republican was nearly 2.1 standard deviations more partisan than the average Republican a decade earlier (2011-2012).¹⁵ This partisan asymmetry mirrors the conclusions of Heseltine (2023), which finds more polarization on social media by Republican members of Congress, as measured by the external news organizations they reference or link to in their tweets and posts.

Notably, however, the trends are different when compared to polarization in roll call voting or in campaign contributions. Columns 3 and 4 of Table 3 shows roll call voting polarization increased for both Democrats and Republicans during this time period, but very modestly, a small fraction of the within-party standard deviation at the beginning of 2010s. Columns 5 and 6 of Table 3 shows slightly more polarization in campaign contribution patterns during this time period, with Democrats polarizing more quickly than Republicans. In none of these regressions

¹⁵In Table J.1 in the Supplemental Materials, we replicate Table J.1 using member fixed effects, allowing us to evaluate how much of the trend in extremity is driven by within-member changes (as opposed to replacement of members with different ones). For Republicans, the Session coefficient estimate is only modestly smaller, suggesting that much of the increase in partisanship is due to members changing the way they communicate over this time period.

does polarization increase as rapidly as it does for Republicans in the rhetoric they use on social media and newsletters over this time period.

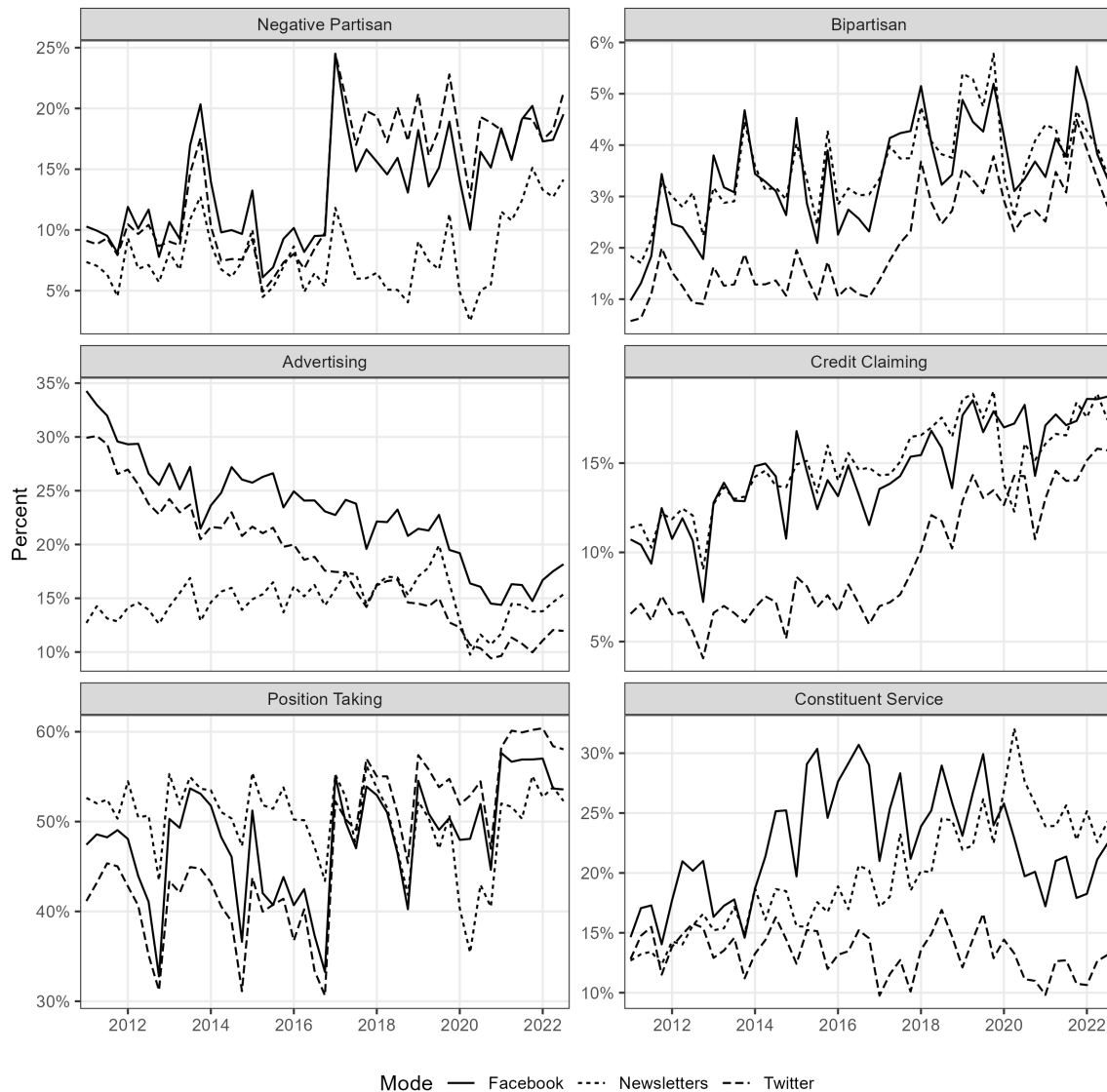
In the Supplemental Materials, we examine the robustness of our rhetorical polarization findings in two ways. First, we consider whether the results replicate if we use a different scaling procedure than the Class Affinity scaling model. In Section K, we re-scale the data using the Naive Bayes approach, which produced the next strongest within-party correlations with the non-text positioning measures such as DW-NOMINATE and CFscores (as discussed in Supplemental Section B). We then estimate identical regression models as in Tables 3 and J.1. The results are shown in Table K.1. When using this measure instead, there is significant evidence of growing partisanship in Democratic rhetoric as well as Republican rhetoric. However, even in these estimates Republican rhetoric polarizes approximately twice as much over the same time period.

Additionally, in Section L we provide a non-model based examination of changes in partisan rhetoric across time by plotting the top 20 words that increased or decreased the most in usage between the 112th and the 117th session of Congress, separately for Democrats and Republicans. As can be seen in Figure L.1, the word changes are suggestive of a shift from routine governance to partisan conflict for both major parties. Democrats increased their emphasis on progressive policy priorities with words like "climate," "gun," "democracy," and "infrastructure" becoming more common, while decreasing references to procedural governance terms like "veteran," "district," "medicare," and "budget." Republicans also show an even more pronounced shift toward combative, identity-focused rhetoric, with large increases in "border," "illegal," "polic," and "china," while dropping more traditional policy language like "regulation," "business," "tax," and "spending." The overall pattern indicates a transformation from governance-focused communication toward more polarized, emotionally charged messaging, with Democrats pivoting to progressive social causes and Republicans to immigration-focused cultural warfare and attacks on the opposing party.

In addition to examining the partisan orientation of speech, we can also examine trends in the six categories of messages displayed in Figure 1. To accomplish this, Figure 5 shows quarterly

averages in the percent of messages classified as Negative Partisan, Bipartisan, Advertising, Credit Claiming, Position Taking, and Constituent Service, separated for each of the three platforms.

FIGURE 5: TRENDS IN PARTISAN AND REPRESENTATIONAL CATEGORIES (2011 - 2022)



Note: The figure displays trends in the quarterly averages of the two partisan and four representational categories trends described above, displayed separately for each type of communication between 2011 and 2022.

Two trends worth remarking on, given their relevance to the partisanship trends discussed already, are sizable increases in both Negative Partisan attacks and Position Taking messages.

Negative Partisan attacks represent approximately 5-10% of messages during the initial years of the time series, but increase to 15-20% by the end of the time period, with a particularly sharp increase towards the end of 2016.¹⁶ Similarly, position taking comprises between 40-55% of messages at the beginning of the time series, but increases to 55-60% by the end of 2022.

Figure M.1 in the Supplemental Materials replicates Figure 5 but displays category trends for each of the two major parties (aggregated across communication platforms). The figure reveals interesting heterogeneity by party. Most notably, the stark increase in Negative Partisan attacks after 2016 is driven by congressional Democrats, most likely in direct response to the election of Republican Donald Trump to the presidency. There is a similar increase in Negative Partisan attacks by Republicans following the election of Democrat Joe Biden. Notably though, even after accounting for factors such as a cross-party president the across time trends still hold. Republicans engage in much higher levels of Negative Partisanship during the Biden presidency compared to the earlier Obama presidency years.

The Social Media Feedback Mirage

In this section, we explore one possible explanation for the increase in partisanship, negativity, and position-taking in online rhetoric during this time period: differential engagement by message type on social media. Prior research has found that incivility or negativity on social media leads to more engagement in more limited data (Ballard et al. 2022; Macdonald, Russell, and Hua 2023; Yu, Wojcieszak, and Casas 2024). We replicate and extend this research, examining engagement as a function of partisan extremity and our six purposive categorizations on both Facebook and Twitter across a 14-year period.

An advantage of the social media data is that posts and tweets contain information about public engagement in the form of likes, shares, retweets, and other metrics. These metrics represent

¹⁶Curiously, bipartisan messages also become more common during this time period. The increase, however, is small in absolute terms, changing by only a few percentage points.

the information members and their staff receive in real time, showing how followers respond to the messages they share. In the case of communications staffers who often run these social media accounts, increasing engagement and the number of followers is often a component of how their performance is evaluated. Given these incentives, it is important to consider which types of messages receive the most attention (particularly positive attention) on social media.

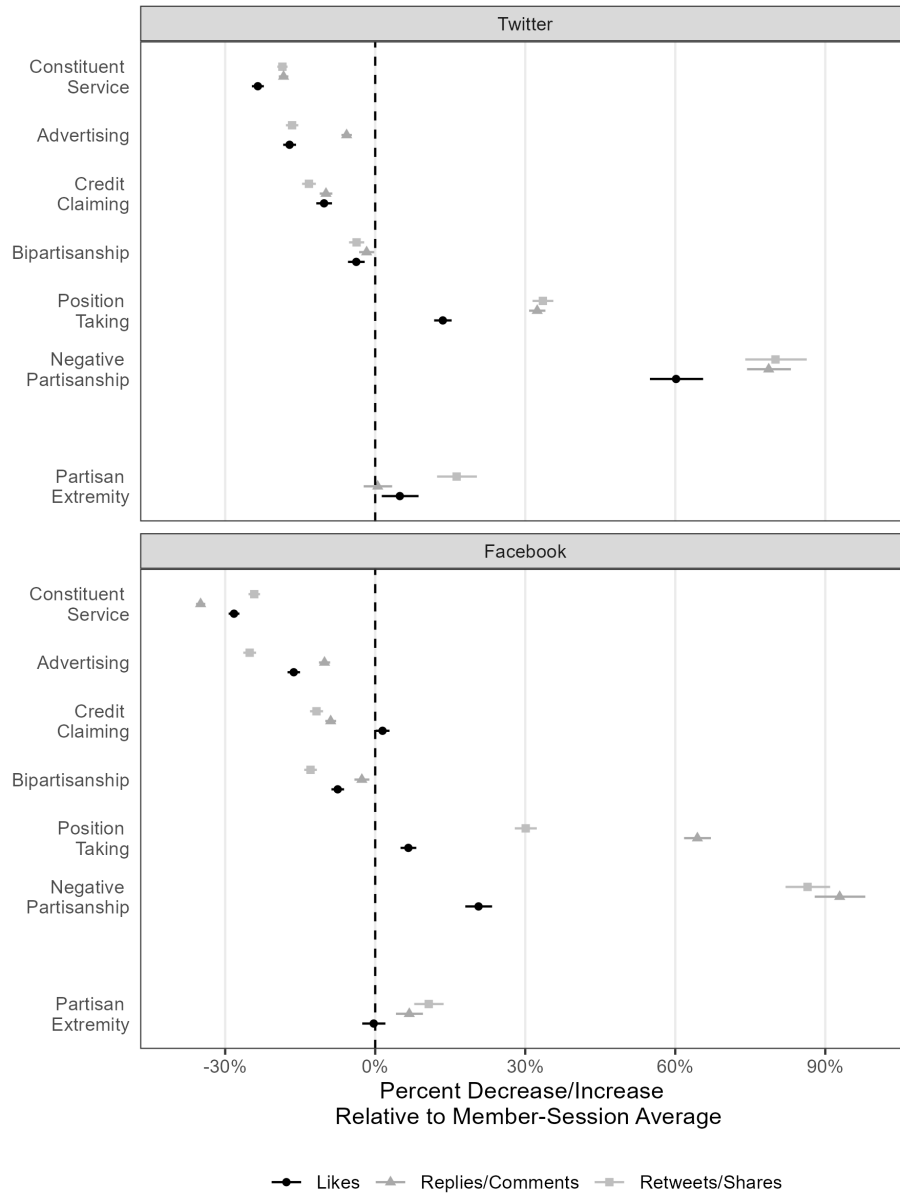
To determine which characteristics of messages in our data receive this attention, we estimate a series of OLS regression models where the unit of observation is an individual tweet or post. For each message, we include six binary variables for whether the message is classified into each of our categories, as well as the continuous (0-1) Text Partisan Extremity Score. The dependent variable is the number of likes, retweets/shares, or replies/comments the message received (logged to address right-skew). Regressions include member-session fixed effects. As a consequence of the fixed effect log-DV specification, coefficients can be interpreted as the percent difference in expected likes or retweets relative to the average tweet or post by a member.

The results are displayed in Figure 6. The figure shows clear differences in engagement across message types. The most striking pattern is the high levels of engagement that negative partisan messages receive. Negative partisan attacks on Twitter receive 60% more likes, 80% more retweets, and 79% more replies than a member's average tweet. Similarly, negative partisan attacks on Facebook receive 21% more likes, 86% more shares, and 93% more comments than a member's average Facebook post.

Generally speaking, there are three types of messages that receive higher levels of engagement: negative messages, more partisan messages, and position-taking. In contrast, the other four types of messages (constituent service, advertising, credit claiming and bipartisanship) receive less engagement than the typical social media message for a particular member.

Does this differential reception by social media users matter? There are at least two reasons it might. One is if members mistake engagement on social media with the preferences of voters and other key constituencies they care about. In general, survey and experimental studies find

FIGURE 6: SOCIAL MEDIA MESSAGE CONTENT AND USER ENGAGEMENT



Note: The figure displays OLS coefficient estimates and 95% confidence intervals for member-session fixed effects models where the dependent variables are the social media engagement metrics shown in the legend at bottom. Independent variables are message characteristics shown on the y-axis. Coefficient estimates are transformed ($\exp(\hat{\beta}) - 1$) to percent differences for interpretability; transformed coefficients thus represent engagement relative to a member's average in a given session of Congress. Standard errors are clustered by member. Full regression results shown in Table I.2 in the Supplemental Materials.

opposite patterns. Respondents report approving more highly of a hypothetical representative they describe their accomplishments in office (credit claiming) versus attacking the other party (e.g., [Barron and McLaughlin 2024](#); [Costa 2021](#); [Simas et al. 2025](#)). Social media thus threatens to provide members a misleading impression about what constituents want out of their representatives if members use social media engagement to provide feedback on what resonates with potential voters. This illusory mismatch between the messaging voters prefer and the messaging that attracts engagement online might be called a “social media feedback mirage”. Whether elected officials draw information from social media engagement in this manner or not is an open question, one that future research should pursue.

A second reason this distortion might matter is if the political discourse of the mass public is influenced by the conversations of political elites on social media. The extent to which public rhetoric follows elite rhetoric and vice-versa is an area of active debate in political science and the broader communication literature (e.g., [Barberá et al. 2019](#); [Smith et al. 2025](#); [Warner 2023](#); [Zoizner and Levy 2025](#)). Despite this, some degree of influence by elites on mass opinion is a longstanding axiom of political science ([Lippmann 1922](#); [Jacobs and Shapiro 2000](#); [Zaller 1992](#)). To the extent that politicians emphasize divisive messages on social media because they garner attention (or equivalently, to the extent that algorithms prioritize such messages), the downstream consequences are likely a shifting of the way members of the public think about and debate politics.

Conclusion

Communications from elected officials both facilitate effective representation (e.g. [Pitkin 1967](#)) and shape the subject and tenor of public conversation. As such, studying the content of those communications is central to numerous strands of the political science literature. Work in this area has attempted to keep pace with the changes that have come with the wide-spread adop-

tion of online communications, but efforts have been hampered by the challenge of collecting such large volumes of data over time and across platforms. Our Scaled and Classified Congressional Communication (SCCC) dataset, spanning approximately 4.7 million tweets, 2.4 million Facebook posts, and 184,000 email newsletters authored by members of Congress between 2009 and 2022, helps advance study in this area in several ways. The size and scope of this dataset not only resolves practical issues related to data collection, but also allows for more comprehensive investigations of how communications change over time and vary across modes. Moreover, the coding of both purpose and partisanship allow for more precise estimates of the effects of the various components of online messages. Altogether, the SCCC is one of the most comprehensive resources available for studying how characteristics of politicians' rhetoric affect a host of key political outcomes.

Indeed, the applications presented here demonstrate the vast potential of this dataset to make contributions beyond just the more focused study of online congressional communications. Our findings that negative and partisan messages are increasingly prevalent and attention-grabbing speak to work on polarization. They also raise interesting questions for further study. Namely, why has online communication grown and evolved to be more partisan in nature when considerable evidence suggests that there are no apparent payoffs from voters? While some of this may be due to disconnect between the feedback received from social media users vs. direct constituents (i.e., the social media feedback mirage), it seems unlikely that savvy politicians would continue to pursue strategies that do not help and possibly hinder their goals in some form. The obvious conclusion is that more work is needed to gain a complete understanding of the motivations behind and consequences of politicians' communications.

It should also not be lost that while a major strength of our approach is that it offers estimates of both style *and* partisanship. The partisan ratings alone are an extremely valuable tool for approximating legislators' positions, in some cases providing more facially valid and/or informative estimates than alternatives such as DW-NOMINATE. For one, the two widely used measures of

legislative positioning we compare our measures to above are derived from legislative voting and fundraising. Though the public has access to records of both, the majority of individuals lack knowledge about their representatives' activities in these areas. Our ratings, in contrast, are derived from visible and easily accessible communications that are crafted to portray a legislator as they want to be seen. So when used in combination with these other types of measures, our partisan scaling has the potential to offer a fuller picture of legislators' preferences and speak to questions of whether those preferences appear relatively consistent when estimated from these different sources.

In addition, there is ample evidence that partisan rhetoric has a significant influence on the thoughts and actions of the public, as partisans can be both swayed by their own party and repelled by the opposite (e.g. [Slothuus and Bisgaard 2021](#); [Pink et al. 2021](#); [Nelson 2004](#)). Having measures of the strength of the partisanship in representatives' messages can thus offer greater insight into the parties' powers to persuade and the downstream consequences of increasingly polarized communications.

Our dataset, as currently constructed, ends in 2022. In the very short time since then, social media usage by politicians has continued to evolve. One of the most notable examples is the rise of new social media platforms that are associated with particular points on the ideological spectrum: Truth Social and Bluesky. But our August 2025 search found that while all sitting members of the House and Senate have an X/Twitter account and 99% have a Facebook account, only 46% have an account on Bluesky and only 20% have an account on Truth Social. Moreover, [Figure N.1](#) in the Appendix shows that almost all members have a link to their Facebook and/or Twitter/X account on the homepage of their official website. And 92% of members had a link or pop-up prompting individuals to signup for their newsletters. Thus, it appears that (1) accounts on alternate platforms are coming in addition to and not in replacement of Facebook and/or Twitter/X; and (2) Facebook, Twitter/X, and newsletters are still the three dominant forms of

communication that members of Congress are pushing their constituents to follow.¹⁷

Even so, it is possible that we will continue to observe even more fractionalization, with politically-oriented media consumers choosing platforms that cater in a focused way to their tastes. On the other hand, even among the population of social media users most do not follow their representatives on social media (McCabe et al. 2023), perhaps limiting the impact of this type of political Balkanization. A clear path forward for scholars of political communication, then, is unpacking the effects these new platforms have on rhetoric and discourse. The large percentages of members with linked YouTube and Instagram accounts suggest that shifting study to more visual communication may be particularly fruitful.

Explorations of more visual forms of communication may also help with the continued evolution of the typology of messages. We take the fact that over 70% of communications can be classified into one of our six categories as evidence that (with slight modifications) Mayhew's and Fenno's classic frameworks still clearly capture the majority of congressional communication. While most of the messages that do not fit into one of our categories lack any real political substance, involving things like holiday wishes or recognition of specific constituent achievements, we acknowledge that there may be messaging styles unique to online communication that simply do not map well onto categories originally intended for different forms of communication. As such, the shareable aspects of our data set will allow scholars to develop their own categorization schemes and further highlight the ways in which the content of messages has and has not changed in accordance with changes in the way those messages are delivered.

So in sum, though politics and political communication has changed drastically in the past half century, the fact remains that “if there is to be congruence between the policy preferences of the represented and the policy decisions of the representatives, however, two-way communication

¹⁷When we look at the consumer side, we also see that emerging platforms still lag behind their more established counterparts. An August 2025 Pew survey of U.S. adults found that 38% report that they regularly get news on Facebook and 12% report that they regularly get news on Twitter/X. In contrast, only 2% report getting news from either Truth Social or Bluesky. See <https://www.pewresearch.org/journalism/fact-sheet/social-media-and-news-fact-sheet/>.

between them is a prerequisite” (Fenno 1978, p.241). We offer a comprehensive resource for the continuing and evolving study of representation in the United States. The possibilities our data unlock are too many to list here; we are confident that scholars from across the discipline will benefit from having such a resource at their disposal, and use it in ways beyond what we can imagine.

References

- Albert, Zachary. 2020. "Partisan Policymaking in the Extended Party Network: The Case of Cap-and-Trade Regulations." *Political Research Quarterly* 73 (2): 476–491.
- Ashworth, Scott. 2012. "Electoral Accountability: Recent Theoretical and Empirical Work." *Annual Review of Political Science* 15 (1): 183–201.
- Ballard, Andrew O., Ryan DeTamble, Spencer Dorsey, Michael Heseltine, and Marcus Johnson. 2023. "Dynamics of Polarizing Rhetoric in Congressional Tweets." *Legislative Studies Quarterly* 48 (1): 105–144.
- Ballard, Andrew, Ryan DeTamble, Spencer Dorsey, Michael Heseltine, and Marcus Johnson. 2022. "Incivility in Congressional Tweets." *American Politics Research* 50 (6): 769–780.
- Ban, Pamela, and Jaclyn Kaslovsky. 2025. "Local Orientation in the U.S. House of Representatives." *American Journal of Political Science* 69 (3): 1082–1098.
- Barber, Michael. 2022. "Comparing Campaign Finance and Vote-Based Measures of Ideology." *The Journal of Politics* 84 (1): 613–619.
- Barberá, Pablo, Andreu Casas, Jonathan Nagler, Patrick J. Egan, Richard Bonneau, John T. Jost, and Joshua A. Tucker. 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.
- Barron, Nathan T., and Peter T. McLaughlin. 2024. "Experimental Evidence of the Benefits and Risks of Credit Claiming and Pork Busting." *Legislative Studies Quarterly* 49 (1): 103–129.
- Besley, Timothy. 2006. *Principled Agents? The Political Economy of Good Government*. Oxford: Oxford University Press.

- Blum, Rachel, Lindsey Cormack, and Kelsey Shoub. 2023. "Conditional Congressional Communication: How Elite Speech Varies Across Medium." *Political Science Research and Methods* 11 (2): 394–401.
- Bonica, Adam. 2014. "Mapping the Ideological Marketplace." *American Journal of Political Science* 58 (2): 367–386.
- Bonica, Adam. 2018. "Inferring Roll-Call Scores from Campaign Contributions Using Supervised Machine Learning." *American Journal of Political Science* 62 (4): 830–848.
URL: <https://onlinelibrary.wiley.com/doi/abs/10.1111/ajps.12376>
- Burke, Edmund. 1774. "Speech to the Electors of Bristol." In *The Works of the Right Honourable Edmund Burke*. Vol. 1 London: Henry G. Bohn pp. 446–448. Published 1854–56.
- Butler, Daniel M., Thad Kousser, and Stan Oklobdzija. 2023. "Do Male and Female Legislators Have Different Twitter Communication Styles?" *State Politics & Policy Quarterly* 23 (2): 117–139.
- Campbell, Angus, Philip E. Converse, Warren E. Miller, and Donald E. Stokes. 1960. *The American Voter*. Chicago: University of Chicago Press.
- Canes-Wrone, Brandice, and Michael R. Kistner. 2023. "Local Newspapers and Ideological Accountability in US House Elections." In *Accountability Reconsidered: Voters, Interests, and Information in US Policymaking*, ed. Charles M. Cameron, Brandice Canes-Wrone, Sanford C. Gordon, and Gregory A. Huber. Cambridge: Cambridge University Press pp. 129 – 149.
- Cohen, Joshua. 1989. "Deliberation and Democratic Legitimacy." In *The Good Polity: Normative Analysis of the State*, ed. Alan Hamlin, and Phillip Petit. New York: Blackwell.
- Converse, Philip E. 1964. "The Nature of Belief Systems in Mass Publics." In *Ideology and Discontent*, ed. David E. Apter. New York: Free Press pp. 206–261.

- Cormack, Lindsey. 2016a. "Extremity in Congress: Communications versus Votes." *Legislative Studies Quarterly* 41 (3): 575–603.
- Cormack, Lindsey. 2016b. "Gender and Vote Revelation Strategy in the United States Congress." *Journal of Gender Studies* 25 (6): 626–640.
- Cormack, Lindsey. 2017. "DCinbox - Capturing Every Congressional Constituent E-newsletter from 2009 Onwards." *The Legislative Scholar* 2 (1): 27–36.
- Costa, Mia. 2021. "Citizen Evaluations of Legislator–Constituent Communication." *British Journal of Political Science* 51 (3): 1324–1331.
- Cowburn, Mike, and Marius Sältzer. 2024. "Partisan Communication in Two-Stage Elections: The Effect of Primaries on Intra-Campaign Positional Shifts in Congressional Elections." *Political Science Research and Methods* pp. 1–20.
- Edwards, III, George C. 1985. "Measuring Presidential Success in Congress: Alternative Approaches." *Journal of Politics* 47 (2): 667–685.
- Evans, Heather K., and Jennifer Hayes Clark. 2016. "'You Tweet Like a Girl!': How Female Candidates Campaign on Twitter." *American Politics Research* 44 (2): 326–352.
- Evans, Heather K., Victoria Cordova, and Savannah Sipole. 2014. "Twitter Style: An Analysis of How House Candidates Used Twitter in Their 2012 Campaigns." *PS: Political Science & Politics* 47 (2): 454–462.
- Fenno, Richard F. 1978. *Home Style: House Members in Their Districts*. Scott, Foresman.
- Fowler, Erika Franklin, Michael M. Franz, Gregory J. Martin, Zachary Peskowitz, and Travis N. Ridout. 2021. "Political Advertising Online and Offline." *American Political Science Review* 115 (1): 130–149.

- Frankel, Laura Lazarus, and D. Sunshine Hillygus. 2017. "Niche Communications in Political Campaigns." In *The Oxford Handbook of Political Communication*, ed. Kate Kenski, and Kathleen Hall Jamieson. Oxford University Press.
- Fu, Shu, and William G. Howell. 2020. "The Behavioral Consequences of Public Appeals: Evidence on Campaign Fundraising from the 2018 Congressional Elections." *Presidential Studies Quarterly* 50 (2): 325–347.
URL: <http://onlinelibrary.wiley.com/doi/abs/10.1111/psq.12645>
- Gainous, Jason, and Kevin M. Wagner. 2013. *Tweeting to Power: The Social Media Revolution in American Politics*. Oxford University Press.
- Gentzkow, Matthew, Jesse M. Shapiro, and Matt Taddy. 2019. "Measuring Group Differences in High-Dimensional Choices: Method and Application to Congressional Speech." *Econometrica* 87 (4): 1307–1340.
URL: <https://www.econometricsociety.org/doi/10.3982/ECTA16566>
- Green, Jon, Kelsey Shoub, Rachel Blum, and Lindsey Cormack. 2024. "Cross-Platform Partisan Positioning in Congressional Speech." *Political Research Quarterly* 77 (3): 653 – 668.
- Greenacre, Michael. 2007. *Correspondence Analysis in Practice*. Chapman & Hall/CRC.
- Grimmer, Justin. 2013. *Representational Style in Congress: What Legislators Say and Why It Matters*. Cambridge University Press.
- Grossmann, Matt, and David A. Hopkins. 2024. *Polarized by Degrees: How the Diploma Divide and the Culture War Transformed American Politics*. New York: Cambridge University Press.
- Habermas, Jürgen. 1968. *The Theory of Communicative Action. Vol. 1. Reason and the Rationalization of Society*. Boston: Beacon Press.

- Hayes, Danny, and Jennifer L. Lawless. 2015. "As Local News Goes, So Goes Citizen Engagement: Media, Knowledge, and Participation in US House Elections." *The Journal of Politics* 77 (2): 447–462.
- Hayes, Danny, and Jennifer L. Lawless. 2017. "The Decline of Local News and Its Effects: New Evidence from Longitudinal Data." *The Journal of Politics* 80 (1): 332–336.
- Hemphill, Libby, Annelise Russell, and Angela M. Schöpke-Gonzalez. 2021. "What Drives U.S. Congressional Members' Policy Attention on Twitter?" *Policy & Internet* 13 (2): 233–256.
- Heseltine, Michael. 2023. "Asymmetric Polarization in Online Media Engagement in the United States Congress." *The International Journal of Press/Politics* 30 (1): 299–325.
URL: <https://doi.org/10.1177/19401612231211800>
- Heseltine, Michael. 2024. "Polarizing Online Elite Rhetoric at the Federal, State, and Local Level During the COVID-19 Pandemic." *American Politics Research* 52 (3): 187–202.
- Heseltine, Michael, and Bernhard Clemm von Hohenberg. 2024. "Large language models as a substitute for human experts in annotating political text." *Research & Politics* 11 (1): 20531680241236239.
URL: <https://doi.org/10.1177/20531680241236239>
- Hilden, David J., and Michael R. Kistner. 2025. "Promoting Conspiracy Theories Strategically." 20 (3): 307–337.
- Hunt, Charles R., and Kristina C. Miler. 2025. "How Modern Lawmakers Advertise Their Legislative Effectiveness to Constituents." *The Journal of Politics* 87 (1): 231–245.
URL: <https://www.journals.uchicago.edu/doi/10.1086/730732>
- Jacobs, Lawrence R, and Robert Y Shapiro. 2000. *Politicians don't pander: Political manipulation and the loss of democratic responsiveness*. University of Chicago Press.

- Jungherr, Andreas, and Ralph Schroeder. 2022. *Digital Transformations of the Public Arena*. Elements in Politics and Communication Cambridge University Press.
- URL:** <https://www.cambridge.org/core/elements/digital-transformations-of-the-public-arena/6E4169B5E1C87B0687190F688AB3866E>
- Kaslovsky, Jaclyn, and Michael R. Kistner. 2025. "Responsive Rhetoric: Evidence from Congressional Redistricting." *Legislative Studies Quarterly* 49 (4).
- Laver, Michael, Kenneth Benoit, and John Garry. 2003. "Extracting Policy Positions from Political Texts Using Words as Data." *American Political Science Review* 97 (2): 311–331.
- Lippmann, Walter. 1922. *Public opinion*. Harcourt, Brace, and Company.
- Local Orientation in the U.S. House of Representatives*. 2025. 69: 779–1195.
- Lowe, Will. 2008. "Understanding Wordscores." *Political Analysis* 16 (4): 356–371.
- Macdonald, Maggie, Annelise Russell, and Whitney Hua. 2023. "Negative Sentiment and Congressional Cue-Taking on Social Media." *PS: Political Science & Politics* 56 (2): 201–206.
- URL:** <https://www.cambridge.org/core/journals/ps-political-science-and-politics/article/negative-sentiment-and-congressional-cuetaking-on-social-media/1E0DF82E947BBE25B35D28AEB203B2B2>
- Macdonald, Maggie, Whitney Hua, and Annelise Russell. 2024. "Constrained Communication and Negativity Bias: Gendered Emotional Appeals on Facebook." *Journal of Women, Politics & Policy* 45 (2): 261–274.
- Mansbridge, Jane. 2003. "Rethinking Representation." *The American Political Science Review* 97 (4): 515–528.
- Mayhew, David R. 1974. *Congress: The Electoral Connection*. Yale University Press.

- McCabe, Stefan, Jon Green, Pranav Goel, and David Lazer. 2023. "Inequalities in Online Representation: Who Follows Their Own Member of Congress on Twitter?" *Journal of Quantitative Description: Digital Media* 3.
- McCarty, Nolan. 2019. *Polarization: What Everyone Needs to Know*. Oxford: Oxford University Press.
- Meisels, Mellissa. 2025. "Candidate Positions, Responsiveness, and Returns to Extremism." *The Journal of Politics* . Forthcoming.
- Mill, John Stuart. 1861. *Considerations on Representative Government*. Cambridge: Cambridge University Press.
- Munger, Kevin. 2023. "Temporal Validity as Meta-Science." *Research & Politics* 10 (3): 1–10.
- Nelson, Thomas E. 2004. "Policy goals, public rhetoric, and political attitudes." *The Journal of Politics* 66 (2): 581–605.
- Nguyen, Dat Quoc, Thanh Vu, and Anh Tuan Nguyen. 2020. BERTweet: A pre-trained language model for English Tweets. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, ed. Qun Liu, and David Schlangen. Online: Association for Computational Linguistics pp. 9–14.
URL: <https://aclanthology.org/2020.emnlp-demos.2/>
- Nokken, Timothy P., and Keith T. Poole. 2004. "Congressional Party Defection in American History." 29 (4): 545–568.
- Payson, Julia, Andreu Casas, Jonathan Nagler, Richard Bonneau, and Joshua A. Tucker. 2022. "Using Social Media Data to Reveal Patterns of Policy Engagement in State Legislatures." *State Politics & Policy Quarterly* 22 (4): 371–395.

- Perry, Patrick O., and Kenneth Benoit. 2017. "Scaling Text with the Class Affinity Model." *arXiv*.
URL: <https://arxiv.org/abs/1710.08963>
- Peterson, Erik. 2021. "Paper Cuts: How Reporting Resources Affect Political News Coverage." *American Journal of Political Science* 65 (2): 443–459.
- Pink, Sophia L, James Chu, James N Druckman, David G Rand, and Robb Willer. 2021. "Elite party cues increase vaccination intentions among Republicans." *Proceedings of the National Academy of Sciences* 118 (32): e2106559118.
- Pitkin, Hanna F. 1967. *The Concept of Representation*. University of California Press.
- Poole, Keith T., and Howard Rosenthal. 2000. *Congress: a Political-Economic History of Roll Call Voting*. Oxford University Press.
- Russell, Annelise. 2018a. "The Politics of Prioritization: Senators' Attention in 140 Characters." *The Forum* 16 (2): 331–356.
- Russell, Annelise. 2018b. "U.S. Senators on Twitter: Asymmetric Party Rhetoric in 140 Characters." *American Politics Research* 46 (4): 695–723.
- Russell, Annelise. 2021a. "Gendered Priorities? Policy Communication in the U.S. Senate." *Congress & the Presidency* 48 (3): 319–342.
- Russell, Annelise. 2021b. *Tweeting is Leading: How Senators Communicate and Represent in the Age of Twitter*. Oxford University Press.
- Scherpereel, John A., Jerry Wohlgemuth, and Audrey Lievens. 2018. "Does Institutional Setting Affect Legislators' Use of Twitter?" *Policy & Internet* 10 (1): 43–60.
- Schroeder, Ralph. 2018. *Social Theory after the Internet: Media, Technology, and Globalization*. UCL Press.
URL: <http://www.jstor.org/stable/10.2307/j.ctt20krxdr>

Simas, Elizabeth N., David Hilden, Michael R. Kistner, and Jamie Wright. 2025. "Credit Claiming and Accountability for Legislative Effectiveness." *Center for Effective Lawmaking Working Paper Series*.

URL: https://thelawmakers.org/wp-content/uploads/2025/01/Credit_Claiming.pdf

Slapin, Jonathan B., and Sven-Oliver Proksch. 2008. "A Scaling Model for Estimating Time-Series Party Positions from Texts." *American Journal of Political Science* 52 (3): 705–722.

URL: <https://www.jstor.org/stable/25193842>

Slothuus, Rune, and Martin Bisgaard. 2021. "How political parties shape public opinion in the real world." *American Journal of Political Science* 65 (4): 896–911.

Smith, Alyssa H., Jon Green, Brooke F. Welles, and David Lazer. 2025. "Emergent structures of attention on social media are driven by amplification and triad transitivity." *PNAS Nexus* 4: pgaf106.

Smith, Sarah A., and Annelise Russell. 2022. "Different Chambers, Divergent Rhetoric: Institutional Differences and Policy Representation on Social Media." *American Politics Research* 50 (6): 792–797.

Straus, Jacob R., Raymond T. Williams, Colleen Shogan, and Matthew Glassman. 2016. "Congressional Social Media Communications: Evaluating Senate Twitter Usage." *Online Information Review* 40 (5): 643–659.

Tausanovitch, Chris, and Christopher Warshaw. 2017. "Estimating Candidates' Political Orientation in a Polarized Congress." *Political Analysis* 25 (2): 167–187.

Tillery, Alvin B. 2021. "Tweeting Racial Representation: How the Congressional Black Caucus Used Twitter in the 113th Congress." *Politics, Groups, and Identities* 9 (2): 219–238.

- Tucker, Joshua, Andrew Guess, Pablo Barbera, Cristian Vaccari, Alexandra Siegel, Sergey Sanovich, Denis Stukal, and Brendan Nyhan. 2018. "Social Media, Political Polarization, and Political Disinformation: A Review of the Scientific Literature." *SSRN Electronic Journal* .
URL: <https://www.ssrn.com/abstract=3144139>
- Vesoulis, Abby. 2021. "Rising GOP Star Madison Cawthorn Peddles a New Trumpism." *Time* .
URL: <https://time.com/5931815/madison-cawthorn-post-trump/>
- Warner, Seth B. 2023. "Analyzing Attention to Scandal on Twitter: Elites Sell What Supporters Buy." *Political Research Quarterly* 76 (2): 841–850.
URL: <https://doi.org/10.1177/10659129221114145>
- Yiannakis, Diana Evans. 1982. "House Members' Communication Styles: Newsletters and Press Releases." *The Journal of Politics* 44 (4): 1049–1071.
- Yu, Xudong, Magdalena Wojcieszak, and Andreu Casas. 2024. "Partisanship on Social Media: In-Party Love Among American Politicians, Greater Engagement with Out-Party Hate Among Ordinary Users." *Political Behavior* 46 (2): 799–824.
- Zaller, John. 1992. *The nature and origins of mass opinion*. Cambridge university press.
- Zoizner, Alon, and Avraham Levy. 2025. "How social media users adopt the toxic behaviors of ingroup and outgroup accounts." *Journal of Computer-Mediated Communication* 30 (6): zmaf018.

Supplementary Materials

Measuring Partisanship and Representation in Online Congressional Communications

Contents

A	Past Research on Online Congressional Communication	SM—2
B	Comparing Alternative Scaling Approaches	SM—3
C	Comparing By-Platform Text Partisanship Scores	SM—7
D	Comparing Perceived Partisanship to Scaling Results	SM—8
E	Examining Evolution in Members’ Rhetoric Across Time	SM—9
F	Keyword Plots for Scaling and Classification Results	SM—11
G	Describing the Supervised Classification Procedure	SM—14
H	Full BERTweet Classification Reports	SM—17
H.1	Twitter	SM—17
H.2	Facebook	SM—19
H.3	Newsletters	SM—21
H.4	AUC, FPR, FNR, and ROC	SM—23
I	Tables of Full Regression Estimates	SM—26
J	Within-Member Polarization Estimates	SM—28
K	Rhetorical Polarization Trends, Naive Bayes Measure	SM—29
L	Changes in Word Usage Across Time Period	SM—30
M	Message Type Trends, by Party	SM—31
N	Social Media and Newsletter Links on House Members’ Homepages	SM—32

A Past Research on Online Congressional Communication

TABLE A.1: SUMMARY OF ONLINE CONGRESSIONAL COMMUNICATION STUDIES

Author(s)	Medium	Chamber(s) & Time Period
Ballard et al. (2022)	Twitter	Both, 2009–2020
Ballard et al. (2023)	Twitter	Both, 2009–2020
Barbera et al. (2019)	Twitter	Both, 2013–2014
Cormack (2016a)	Newsletters	Both, 2009–2010
Cormack (2016b)	Newsletters	Both, 2009–2010
Cowburn & Saltzer (2024)	Twitter	House, 2020
Davis & Russell (2024)	Twitter	Senate, 2013–2023
Evans & Clark (2016)	Twitter	House, 2012
Evans et al. (2014)	Twitter	House, 2012
Fowler et al. (2021)	Facebook Ad Library (API)	Both, 2018
Fu & Howell (2020)	Twitter and Facebook	Both, 2018
Greene (2024)	Facebook	Both, 2016–2022
Green et al. (2024)	Twitter, Facebook, newsletters, press releases, and floor speeches	House, 2019–2021
Hemphill et al. (2020)	Twitter	Both, 2017–2019
Heseltine (2023)	Twitter and Facebook	Both, 2011–2022
Heseltine (2024)	Twitter and Facebook	Both, 2020–2022
Hunt & Miler (2025)	Newsletters	House, 2009–2020
Kaslovsky & Kistner (2024)	Twitter	House, 2019–2022
LaPlant et al. (2023)	Twitter	House, 2020
Macdonald et al. (2023)	Facebook	House, 2019–2020
Macdonald et al. (2024)	Facebook	House, 2019–2020
Macdonald et al. (2025)	Twitter	House, 2019–2020
McKee et al. (2021)	Twitter	Both, 2019
Russell (2018a)	Twitter	Senate, 2013, 2015
Russell (2018b)	Twitter	Senate, 2013, 2015
Russell (2021a)	Twitter	Senate, 2015
Smith & Russell (2022)	Twitter	Both, 2017–2019
Straus et al. (2016)	Twitter	Senate, 2014
Tillery (2018)	Twitter	Both, 2013–2014
Warner (2023)	Twitter	Both, 2013–2014
Yu et al. (2024)	Twitter	Both, 2016–2020

B Comparing Alternative Scaling Approaches

Given our focus on estimating partisanship rather than ideology, the primary scaling algorithm we use is the class affinity model (Perry and Benoit 2017). The class affinity model is a form of supervised learning in that it takes as an input the party identification of the speaker, then uses that to predict affinity towards the different parties both within and across texts.

To justify this choice empirically, we compare the results of using the class affinity model to several other commonly used text scaling algorithms. Specifically, we consider *Wordscore* (Laver, Benoit, and Garry 2003), *Wordfish* (Slapin and Proksch 2008), *Naive Bayes*, and *Correspondence Analysis*.

The first alternative method we consider is the *Wordscore* text scaling method. *Wordscores* (Laver, Benoit, and Garry 2003) is a supervised text scaling method that positions political texts by comparing word frequencies in unscored texts to those in reference texts with known positions. The technique calculates scores for individual words based on their relative frequencies in the reference texts, then uses these word scores to estimate positions of unscored documents by taking the weighted average of the scores of words they contain. For most direct comparison with the Class Affinity Scaling model, the known positions of texts are coded as -1 if the speaker is a Democrat, and 1 if the speaker is a Republican, making this approximate a continuous measure of partisanship.

The second alternative method we consider is *Wordfish*, a scaling method developed by Slapin and Proksch (2008). *Wordfish* is a statistical model for analyzing political text that estimates policy positions of political actors (like parties or politicians) based on word frequencies in documents. Unlike *Wordscores*, *Wordfish* is an unsupervised scaling method that doesn't require reference texts or human coding that simultaneously estimates both document positions and word weights. The *Wordfish* algorithm, we believe, is a poor choice for at least two reasons. First, it is described as capturing ideology, rather than partisanship. Second, the computational time of *Wordfish*

scales exponentially with corpus size, meaning that it is computationally infeasible for a corpus of our size. Nevertheless, in our analysis below we apply the method to a sufficiently small stratified random sample that the algorithm is capable of running. Specifically, we estimate the Wordfish model using 500 randomly selected messages from each member of Congress in each session.¹⁸

The third alternative method we consider is Correspondence Analysis. Correspondence Analysis (Greenacre 2007) is an unsupervised scaling method that recovers dimensionality by decomposing a transformed matrix of χ^2 distances. Because of its computational efficiency with sparse matrices, CA is used by Bonica (2014) to construct CFscores. When applied to text data, CA treats the matrix of word frequencies across documents as a two-way frequency table, with rows representing individual texts and columns representing word features. The method is attractive as a scaling strategy because it offers a close approximation of a statistical ideal point model at a much-reduced computational cost compared to likelihood-based approaches. As noted by Lowe (2008), CA is nearly equivalent to a log-linear ideal point model. Moreover, the χ^2 distance metric that CA employs normalizes the document-feature matrix by reweighting rows and columns that are more densely populated than others, which serves a function similar to including document and word fixed effects in a regression framework. This normalization property helps account for systematic differences in document length and word frequency, making CA particularly well-suited for scaling political texts where such variation is common. As with Wordfish, this is best thought of as a measure of text ideology rather than text partisanship.

The fourth and final alternative method we consider is Naive Bayes. Unlike the other methods discussed above, Naive Bayes is fundamentally a supervised classification algorithm rather than a scaling method. The Naive Bayes classifier predicts party membership by calculating the posterior probability that a document belongs to each party class based on the observed word frequencies, under the assumption of conditional independence between features given the class

¹⁸We exclude members with fewer than 500 messages.

label. While this approach does not directly estimate positions on a latent dimension, the predicted probabilities of party membership can serve as a rough approximation of partisan orientation. Specifically, we use the predicted probability that a text was written by a Republican (versus a Democrat) as our measure of partisanship. This probability can be interpreted as a continuous measure of the degree to which a text exhibits partisan language patterns, with values closer to 0 or 1 indicating stronger Democratic or Republican orientation, respectively. We estimate the Naive Bayes model using the party labels of the speakers as the outcome variable and the document-feature matrix as the predictors.

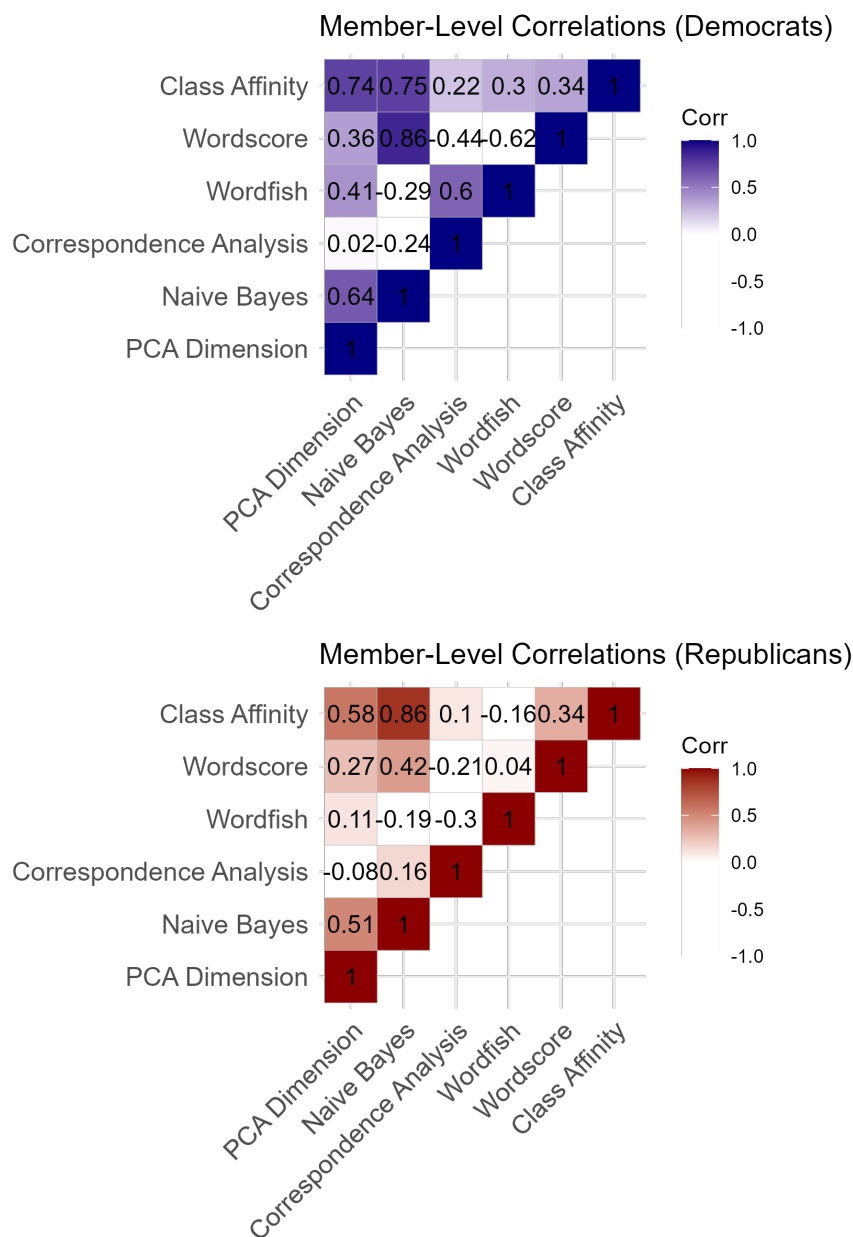
To evaluate whether either of these alternative scaling procedures is more correlated with other member-level ideology and partisanship measures, we conduct a similar exercise to what we show in Figure 2. We first apply the Principal Components Analysis dimension reduction to the pre-existing measures of ideology and partisanship described in that section based on roll call voting, campaign contributions, presidential support and website text. Then, after applying the five text scaling alternatives (Class Affinity, Wordscores, Wordfish, Correspondence Analysis, and Naive Bayes) to the stratified random sample of messages across the three platforms, we obtain a member average. These member averages for each scaling method are then correlated with each other as well as the PCA Dimension, again done separately for each party.

The results are shown in Figure B.1. As the figure demonstrates, the Class Affinity Scaling Model has the highest within-party correlation with the PCA dimension for both Democrats and Republicans. The next best performing algorithm – in terms of correlation with the PCA dimension – is the Naive Bayes approach. There is a strong correlation between the NB predicted probabilities and both the PCA dimension and the Class Affinity results. While we believe the NB approach is a defensible one, the slightly weaker correlation with the PCA composite of the other positioning measures, the fact that it is fundamentally not a scaling algorithm, and the interpretability of the Class Affinity approach make us prefer the latter.

Turning to the other measures, Wordscore also has a modestly strong correlation with the

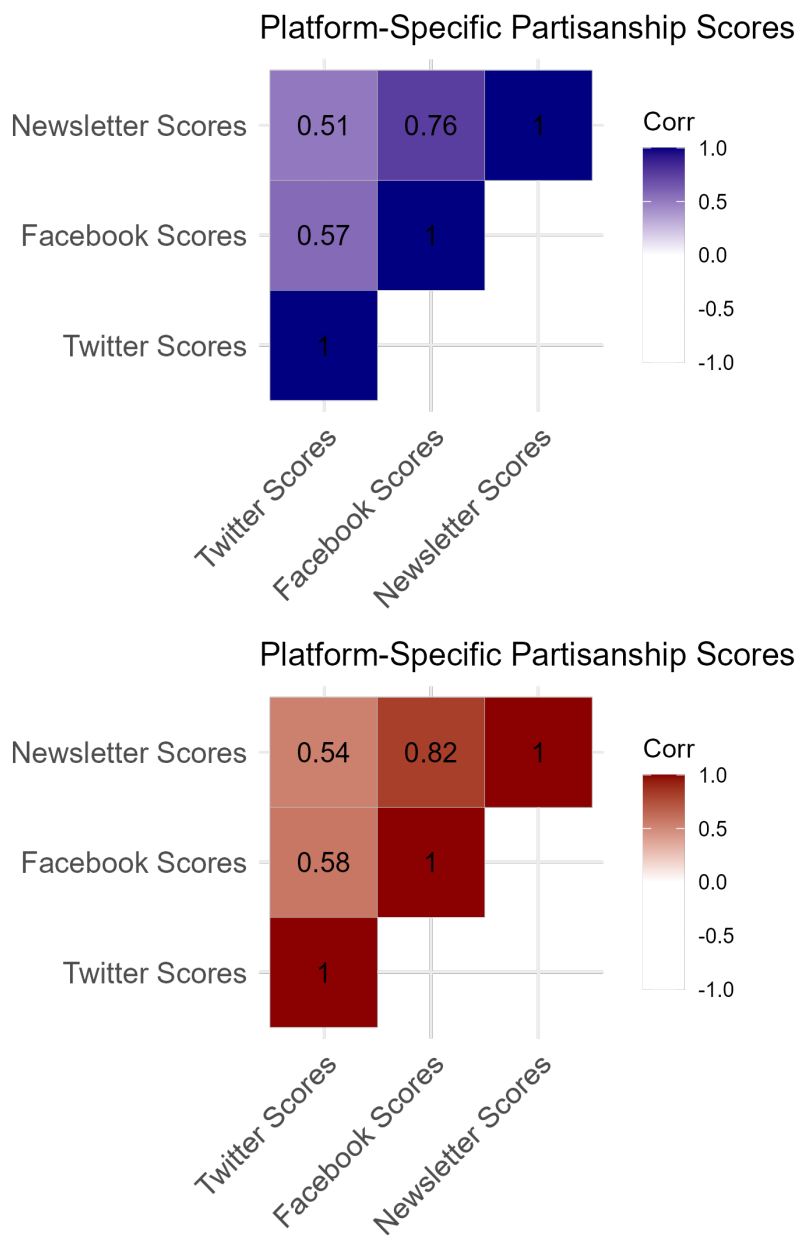
PCA dimension for members of both parties. Wordfish is correlated fairly high with the PCA dimension for Democrats, but not for Republicans. Correspondence Analysis performs poorly in both cases. We view these results as validating our choice of scaling algorithm based on performance reasons.

FIGURE B.1: CORRELATION OF TEXT SCALING MODELS AND LATENT POSITIONING



C Comparing By-Platform Text Partisanship Scores

FIGURE C.1: COMPARING BY-PLATFORM TEXT PARTISANSHIP SCORES

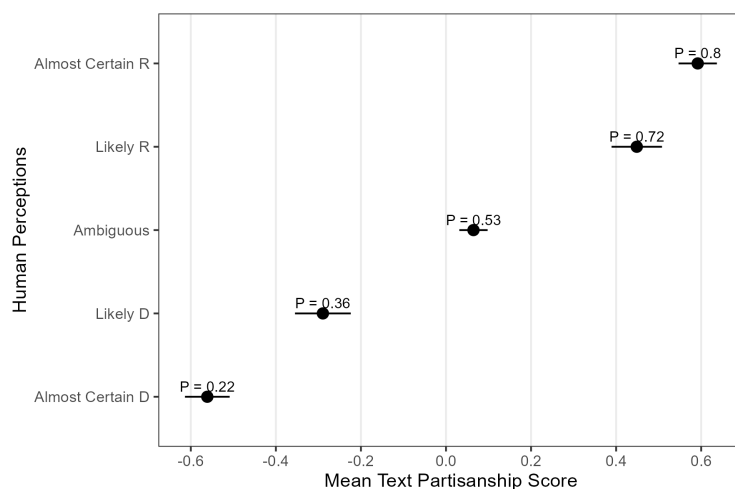


Note: The figure displays the within-party member-level correlations between text partisanship scores constructed separately for each of the three platforms we study. Numbers and shading of each cell shows the Pearson correlation.

D Comparing Perceived Partisanship to Scaling Results

Here we take advantage of the fact that our Text Partisanship Scores are designed to capture the probability that the speaker is Republican (versus Democratic). To ensure that our constructed scores correspond to actual human perceptions, we provided a set of coders 3,000 anonymized messages and asked them to classify the message as originating Almost Certainly from a Republican (Democrat), Likely from a Republican (Democrat), or Ambiguous in speaker partisanship. As Figure D.1 demonstrates, the Text Partisanship Scores of messages align very closely with human perceptions. Messages humans coded as Almost Certainly Republican (Democratic) have a mean probability of 0.80 (0.78) of coming from a Republican (Democrat), while messages humans coded as Likely Republican (Democratic) have a mean probability of 0.72 (0.64) of coming from a Republican (Democrat). Messages coded as Ambiguous have a mean probability of 0.53 of coming from a Republican. Messages coded as Likely Democratic have a mean probability of 0.36 of coming from a Republican. Messages coded as Almost Certainly Democratic have a mean probability of 0.22 of coming from a Republican.

FIGURE D.1: COMPARING TEXT PARTISANSHIP SCORES TO HUMAN PERCEPTIONS OF SPEAKER PARTISAN IDENTITY



Note: The figure displays the mean Text Partisanship Score for anonymized messages classified by human coders into one of five categories based on the perceived partisan identity of the speaker. Solid lines indicated 95% confidence intervals. Annotations above each point show the corresponding probability (π_r) the speaker is Republican.

E Examining Evolution in Members' Rhetoric Across Time

In this Appendix Section, we consider three prominent cases where member partisanship was widely perceived to have changed in a relatively short time period.

We first consider the evolving partisanship of Justin Amash, the former representative from Michigan, who in 2019 changed his party affiliation from Republican to Independent amidst growing public opposition to Republican President Trump in his first term.¹⁹

His movement from more to less Republican than his average counterpart is consistent with his increasing defection from the party and public opposition to Trump.

Second, we consider the evolving partisanship of Senator Lindsey Graham (R-SC), who went from being one of Trump's most outspoken critics leading up to and during the 2016 election to becoming one of Trump's most loyal allies by the second half of his first term.²⁰

Third, we consider the evolving partisanship of Joe Manchin, former senator from West Virginia. Though consistently less partisan than his fellow Democrats in the Senate, Manchin was publicly criticized in particular for his support of Trump and his nominees shortly after Trump first became president.²¹ He then has subsequent falling-out with Trump after his impeachment vote, followed by a return to more Democratic opposition under Biden.

Figure E.1 displays the by-session Text Partisanship scores for all three of these members of Congress. In Figure E.1, all scalings are standardized so that positive (negative) values indicate that the individual was more (less) partisan than the average legislator from their party and chamber in a given session.

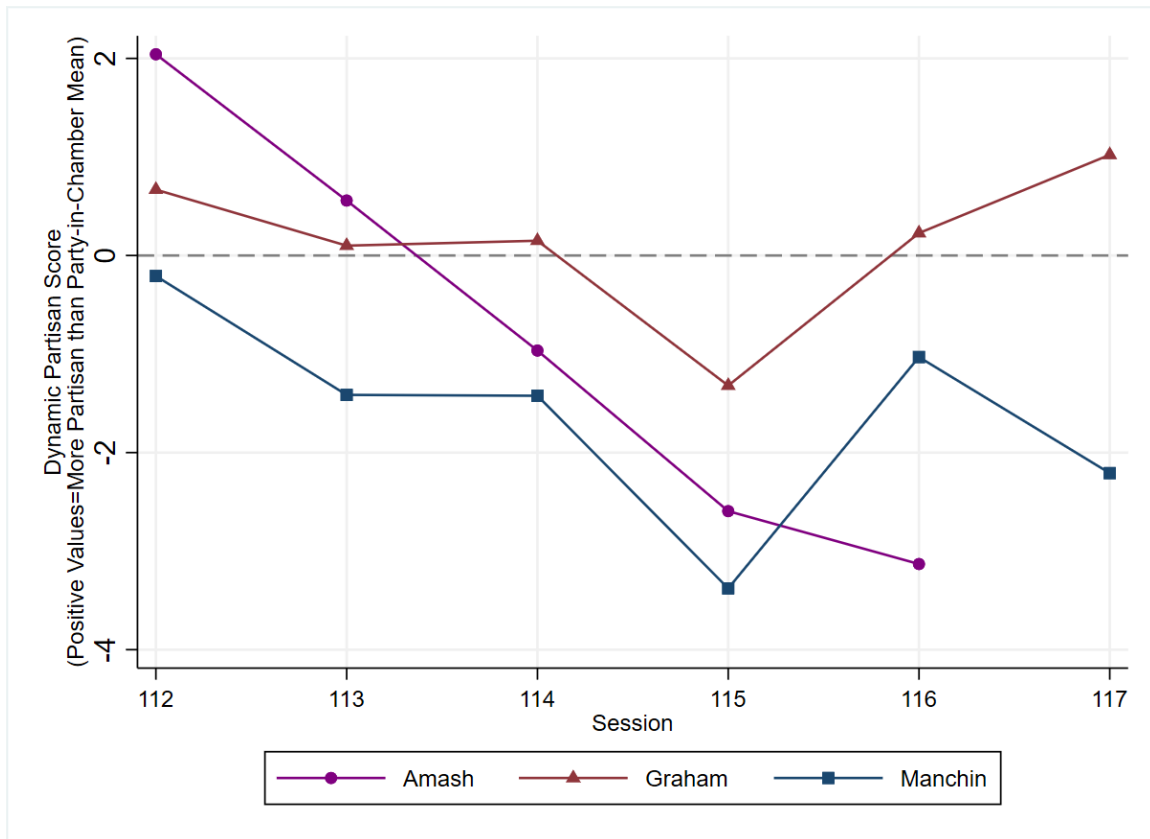
The evolving Text Partisan Scores of each member reflect the dynamics described above. Amash begins as more Republican than the typical member of his party, but by his final term

¹⁹For more details, see <https://michiganadvance.com/2019/06/12/amash-votes-against-his-party-more-than-any-other-u-s-house-member/>

²⁰For more details, see <https://www.nytimes.com/2019/02/25/magazine/lindsey-graham-what-happened-trump.html>

²¹For more details, see <https://www.cnn.com/2017/04/27/politics/joe-manchin-trump-hundred-days>

FIGURE E.1: THE CHANGING PARTISAN RHETORIC OF AMASH, GRAHAM, AND MANCHIN



in office is considerably less Republican than the remaining Republicans in Congress. Graham in the pre-Trump years is quite similar in Text Partisanship to other Republicans, but becomes less Republican in the first two years of Trump's first term (the 115th Congress). Over the next four years, Graham quickly becomes more partisan than the typical Republican. Finally, Manchin moves from being slightly less partisan than the typical Democrat to becoming much less partisan in the 115th Congress, followed by a reversion to most Democrats around the time of Trump's first impeachment, and a further distancing from Democrats in Biden's first two years in office.

F Keyword Plots for Scaling and Classification Results

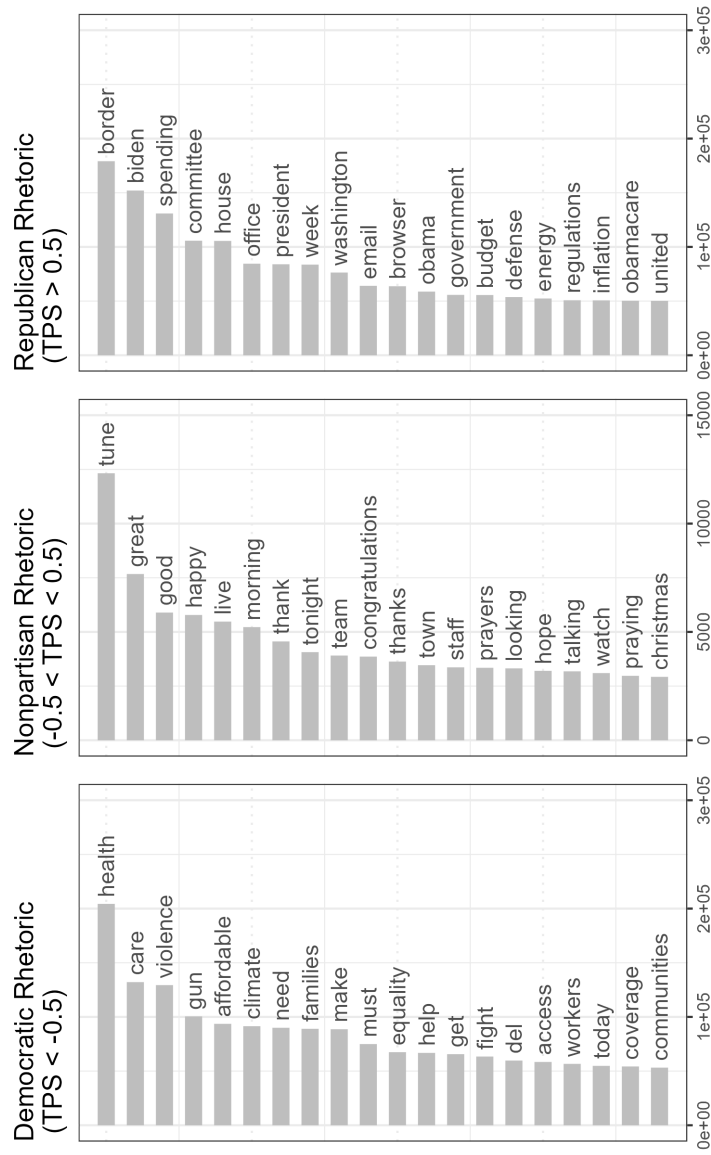
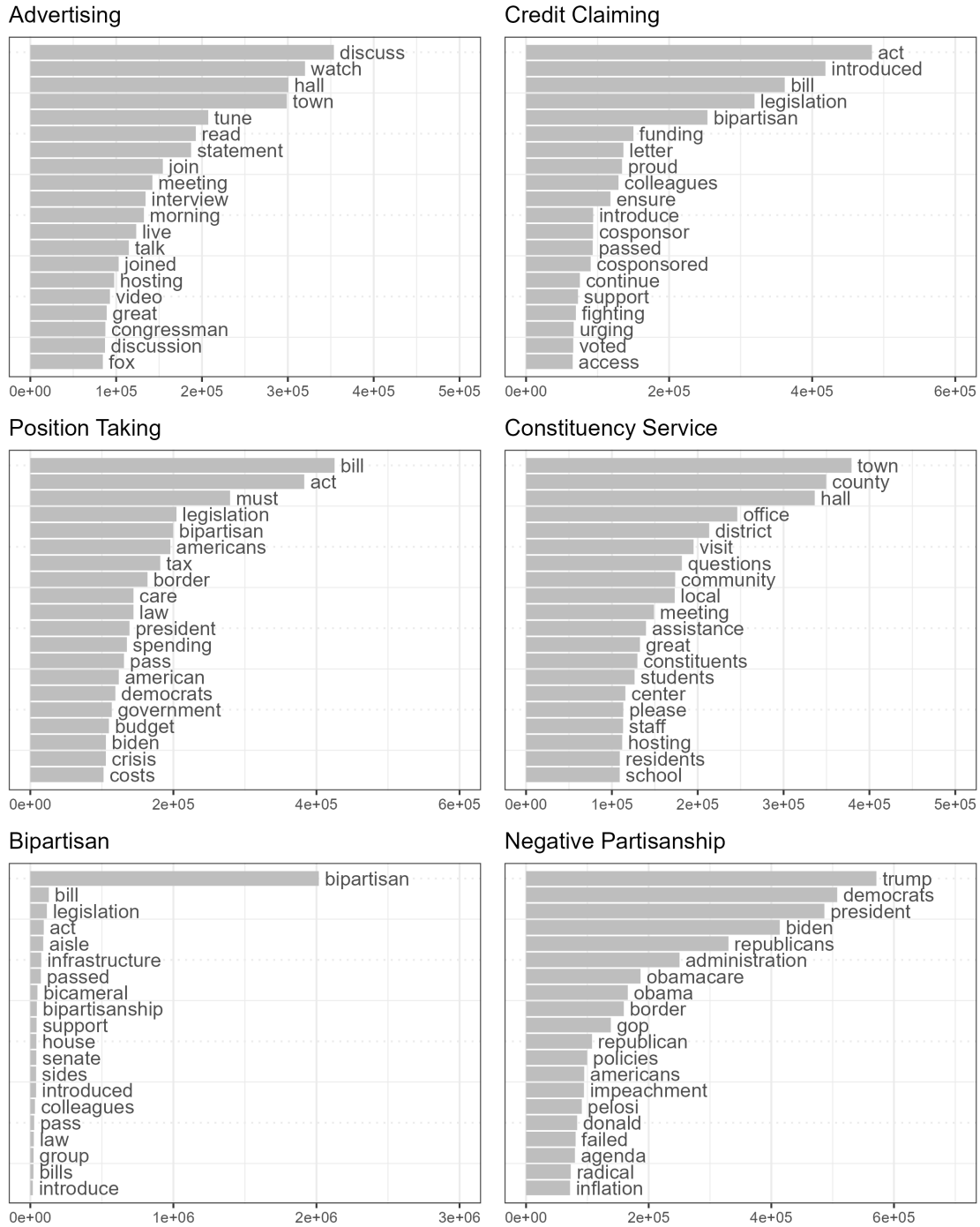


FIGURE F.1: MOST PREDICTIVE WORDS FOR TEXT PARTISANSHIP SCALE

Note: Frequency of key terms for scaling at various points. The left panel shows Democratic rhetoric (a Text Partisanship Score of -1 to -1/2), the middle panel shows mixed rhetoric (a Text Partisanship Score of -1/2 to +1/2), and the right panel shows Republican rhetoric (a Text Partisanship Score of +1/2 to 1).

FIGURE F.2: MOST PREDICTIVE WORDS FOR CATEGORIES



Note: Figure displays the 20 words most strongly associated with classification into each of the six main categories. Likelihood ratio shown along x-axis.

G Describing the Supervised Classification Procedure

First, members of the research team classified a stratified random sample of tweets posted by members of Congress across the entire time span (2009 to 2022), sampled to ensure an approximately equal number of tweets for each chamber-year combination. 8,100 of these tweets were read separately by three members of the research team, to allow us to assess the intercoder reliability. In all categories, the Krippendorff’s alpha was 0.90 or above, indicating high levels of intercoder reliability.²²

The coding team then classified an additional 4,000 Facebook posts and 4,000 newsletters sentence bigrams, to ensure that a substantial amount of platform-specific data was included in our combined training data. We then withheld 1,500 rows from each platform type (4,500 in total, 1,500 tweets, 1,500 Facebook posts, and 1,500 newsletter bigrams) as out-of-sample validation data and used the remaining 16,500 rows for training.²³

A variety of classification algorithms were tested, with out-of-sample balanced accuracy and F1 score used to choose among alternatives. Ultimately, the BERTweet language model (Nguyen, Vu, and Tuan Nguyen 2020), a RoBERTa variant pre-trained specifically on English-language tweets, proved to have the best performance.²⁴ Besides being pre-trained on social media data,

²²The Krippendorff’s alpha for each category was: Advertising - 0.93, Bipartisan - 0.96, Credit Claiming - 0.93, Position Taking 0.94, Negative Partisan - 0.97, and Constituent Service - 0.96 (the constituent service category was created after the initial round of coding and is therefore only coded by two coders with a cross-over set of 500 tweets). We also have coders classify credit claiming into two separate categories: credit claiming for distributive goods (funding and projects for members’ districts and states) and credit claiming for policy work (sponsoring and passing legislation not targeted at a member’s specific constituency). The intercoder reliability for the distributive credit claiming category was 0.90, while for the policy credit claiming category was 0.92.

²³We also included a slight oversample for the Bipartisan category, as this had the lowest positive incidence rate, especially on Twitter. To do this, we used LLaMa 3.3 to classify an additional sample of tweets using a simple prompt to label bipartisan messages. We then pulled the first 100 messages labeled as bipartisan, thus giving us a small additional set of likely bipartisan messages (while avoiding simple keyword filters). We then manually coded this additional sample and added 80 to our training data and 20 to our Twitter validation set.

²⁴An alternative to the supervised learning procedure described here would be to use a large language model to perform classifications, either using one-shot learning or fine-tuning a model. Such an approach is less than optimal here, for multiple reasons. First, using a proprietary LLM for such a task would be prohibitively expensive given the approximately 10 million messages that would need to be classified. Second, the classifications themselves would be subject to change based on updates to the LLM, which occur regularly, producing instability and limiting

the BERTweet classifier also uses vector embeddings to read messages in their entirety with the original ordering preserved, as opposed to more simplistic bag-of-words approaches, making it unsurprising that the accuracy outperformed alternatives such as Naive Bayes and random forests models. The training data were used to tune the base BERTweet model for three epochs, with a learning rate of $2e-5$.

Out-of-sample classification performance metrics for all platforms (accuracy, balanced accuracy, and macro F1) are displayed in Table G.1, while more detailed classification metrics are provided in section H of the Supplemental Materials.²⁵ The macro F1 for each category of tweets ranged from 0.83 to 0.95, values indicating strong performance. Classifier performance on the Facebook data was substantively identical. with Macro F1 scores ranging from 0.81 to 0.92. The models showed marginally weaker performance on the newsletter data, reflecting the more freeform nature of the text compared to social media, but still well within acceptable levels (Macro F1 between 0.77 and 0.93), demonstrating the possibility of cross-platform application of our classification models. For comparison’s sake, the bottom of Table G.1 displays the accuracy, balanced accuracy, and F1 score (as reported) for six recently published research articles classifying social media messages by members of Congress (Ballard et al. 2022, 2023; Yu, Wojcieszak, and Casas 2024), U.S. state legislators (Butler, Kousser, and Oklobdzija 2023; Payson et al. 2022), or both (Fowler et al. 2021). As can be seen, our models show comparatively strong performance across categories and across text types.²⁶

replicability. Despite this, we assessed the relative classification success on a small subset (2,000 tweets) of our data using a fine-tuned GPT-4o mini model. Differences in classification were minimal. As also shown by Heseltine and von Hohenberg (2024), the difference in downstream results based on LLM or transformer-based classification is negligible, at a fraction of the expense or computational resource requirements.

²⁵For evaluating out-of-sample performance for Facebook and newsletters, we use sentence bigrams because they are approximately the same length as tweets, on which the model was primarily pre-trained and trained. Approximately half (45.5%) of Facebook posts are two sentences or fewer and two-thirds (69%) are three sentences or fewer, when applying our models to make final classifications, the Facebook models are applied to entire posts. Because newsletters are considerably longer, final classifications are made on sentence bigrams. Unless otherwise noted, analyses for newsletters are conducted at the sentence bigram level.

²⁶While not displayed in the results of Table G.1 below, the out-of-sample metrics for the two distinct credit claiming categories were also moderately strong. The credit claiming for distributive goods classifier had an out-of-sample

TABLE G.1: CLASSIFICATION ACCURACY METRICS BY CATEGORY AND MODE, WITH COMPARISONS

Out-of-Sample Accuracy Metrics				
Platform	Category	Accuracy	Balanced Accuracy	F1
Twitter	Advertising	0.88	0.84	0.83
	Credit Claiming	0.94	0.83	0.84
	Position Taking	0.88	0.88	0.88
	Constituent Service	0.96	0.93	0.90
	Bipartisanship	0.99	0.94	0.95
	Negative Partisanship	0.95	0.91	0.91
	Range:	0.88–0.99	0.83–0.94	0.83–0.95
Facebook	Advertising	0.88	0.82	0.81
	Credit Claiming	0.93	0.85	0.86
	Position Taking	0.88	0.88	0.88
	Constituent Service	0.96	0.96	0.92
	Bipartisanship	0.99	0.90	0.90
	Negative Partisanship	0.96	0.91	0.88
	Range:	0.88–0.99	0.82–0.96	0.81–0.92
Newsletters	Advertising	0.89	0.78	0.77
	Credit Claiming	0.90	0.82	0.85
	Position Taking	0.84	0.84	0.84
	Constituent Service	0.92	0.90	0.89
	Bipartisanship	0.99	0.90	0.93
	Negative Partisanship	0.96	0.88	0.82
	Range:	0.84–0.99	0.78–0.90	0.77–0.93
Comparisons From Other Published Work				
Article	Platform	Accuracy	Balanced Accuracy	F1
Fowler et al. (2021)	Facebook	0.80–0.99	0.50–0.95	–
Payson et al. (2022)	Tweets	0.55–0.59	–	0.35–0.92
Ballard et al. (2022)	Tweets	0.63–0.97	–	0.64–0.97
Ballard et al. (2023)	Tweets	–	–	0.75–0.94
Butler et al. (2023)	Tweets	0.20–0.99	–	–
Yu et al. (2024)	Tweets	0.66–0.74	0.67–0.85	0.50–0.77
	All:	0.20–0.99	0.50–0.95	0.35–0.97

Note: The top portion of the table displays the accuracy, balanced accuracy, and F1 score (out-of-sample) for each of the six representational categories for tweet, Facebook post, and newsletter sentence data. The bottom portion of the table displays the equivalent metrics reported in recently published research articles using supervised classification techniques to classify social media messages by legislators, to give context for classifier performance.

accuracy of 0.94, 0.91, and 0.80 for the tweet, Facebook, and newsletter validation sets (respectively). The credit claiming for policy work classifier had an out-of-sample accuracy of 0.94, 0.94, and 0.90 for same validation sets.

H Full BERTweet Classification Reports

H.1 Twitter

TABLE H.1: ADVERTISING CLASSIFIER PERFORMANCE, TWITTER VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.941	0.906	0.923	1191
1	0.683	0.780	0.728	309
Accuracy			0.880	1500
Macro avg	0.812	0.843	0.826	1500
Weighted avg	0.888	0.880	0.883	1500

TABLE H.2: BIPARTISAN CLASSIFIER PERFORMANCE, TWITTER VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.995	0.997	0.996	1440
1	0.914	0.883	0.898	60
Accuracy			0.992	1500
Macro avg	0.954	0.940	0.947	1500
Weighted avg	0.992	0.992	0.992	1500

TABLE H.3: CREDIT CLAIMING CLASSIFIER PERFORMANCE, TWITTER VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.961	0.974	0.967	1334
1	0.764	0.681	0.720	166
Accuracy			0.941	1500
Macro avg	0.862	0.827	0.843	1500
Weighted avg	0.939	0.941	0.940	1500

TABLE H.4: CONSTITUENT SERVICE CLASSIFIER PERFORMANCE, TWITTER VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.986	0.968	0.977	1335
1	0.772	0.885	0.825	165
Accuracy			0.959	1500
Macro avg	0.879	0.926	0.901	1500
Weighted avg	0.962	0.959	0.960	1500

TABLE H.5: NEGATIVE PARTISAN CLASSIFIER PERFORMANCE, TWITTER VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.970	0.971	0.970	1250
1	0.855	0.848	0.851	250
Accuracy			0.951	1500
Macro avg	0.912	0.910	0.911	1500
Weighted avg	0.951	0.951	0.951	1500

TABLE H.6: POSITION TAKING CLASSIFIER PERFORMANCE, TWITTER VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.884	0.868	0.876	726
1	0.878	0.893	0.885	774
Accuracy			0.881	1500
Macro avg	0.881	0.880	0.880	1500
Weighted avg	0.881	0.881	0.881	1500

H.2 Facebook

TABLE H.7: ADVERTISING CLASSIFIER PERFORMANCE, FACEBOOK VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.936	0.916	0.925	1221
1	0.662	0.724	0.692	279
Accuracy			0.880	1500
Macro avg	0.799	0.820	0.809	1500
Weighted avg	0.885	0.880	0.882	1500

TABLE H.8: BIPARTISAN CLASSIFIER PERFORMANCE, FACEBOOK VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.993	0.994	0.994	1451
1	0.830	0.796	0.812	49
Accuracy			0.988	1500
Macro avg	0.911	0.895	0.903	1500
Weighted avg	0.988	0.988	0.988	1500

TABLE H.9: CREDIT CLAIMING CLASSIFIER PERFORMANCE, FACEBOOK VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.958	0.965	0.961	1300
1	0.759	0.725	0.742	200
Accuracy			0.933	1500
Macro avg	0.859	0.845	0.851	1500
Weighted avg	0.931	0.933	0.932	1500

TABLE H.10: CONSTITUENT SERVICE CLASSIFIER PERFORMANCE, FACEBOOK VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.993	0.957	0.974	1266
1	0.804	0.962	0.875	234
Accuracy			0.957	1500
Macro avg	0.898	0.959	0.925	1500
Weighted avg	0.963	0.957	0.959	1500

TABLE H.11: NEGATIVE PARTISAN CLASSIFIER PERFORMANCE, FACEBOOK VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.985	0.965	0.975	1360
1	0.719	0.857	0.782	140
Accuracy			0.955	1500
Macro avg	0.852	0.911	0.878	1500
Weighted avg	0.960	0.955	0.957	1500

TABLE H.12: POSITION TAKING CLASSIFIER PERFORMANCE, FACEBOOK VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.895	0.873	0.884	800
1	0.858	0.883	0.870	700
Accuracy			0.877	1500
Macro avg	0.877	0.878	0.877	1500
Weighted avg	0.878	0.877	0.877	1500

H.3 Newsletters

TABLE H.13: ADVERTISING CLASSIFIER PERFORMANCE, NEWSLETTER VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.958	0.921	0.939	1348
1	0.480	0.645	0.551	152
Accuracy			0.893	1500
Macro avg	0.719	0.783	0.745	1500
Weighted avg	0.910	0.893	0.900	1500

TABLE H.14: BIPARTISAN CLASSIFIER PERFORMANCE, NEWSLETTER VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.992	0.997	0.994	1437
1	0.911	0.810	0.857	63
Accuracy			0.989	1500
Macro avg	0.951	0.903	0.926	1500
Weighted avg	0.988	0.989	0.988	1500

TABLE H.15: CREDIT CLAIMING CLASSIFIER PERFORMANCE, NEWSLETTER VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.903	0.980	0.940	1146
1	0.910	0.658	0.764	354
Accuracy			0.904	1500
Macro avg	0.906	0.819	0.852	1500
Weighted avg	0.904	0.904	0.898	1500

TABLE H.16: CONSTITUENT SERVICE CLASSIFIER PERFORMANCE, NEWSLETTER VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.957	0.929	0.943	1132
1	0.800	0.872	0.835	368
Accuracy			0.915	1500
Macro avg	0.879	0.901	0.889	1500
Weighted avg	0.919	0.915	0.917	1500

TABLE H.17: NEGATIVE PARTISAN CLASSIFIER PERFORMANCE, NEWSLETTER VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.989	0.972	0.981	1430
1	0.579	0.786	0.667	70
Accuracy			0.963	1500
Macro avg	0.784	0.879	0.824	1500
Weighted avg	0.970	0.963	0.966	1500

TABLE H.18: POSITION TAKING CLASSIFIER PERFORMANCE, NEWSLETTER VALIDATION SET

	Precision	Recall	F1-score	Support
0	0.887	0.835	0.860	874
1	0.787	0.851	0.818	626
Accuracy			0.842	1500
Macro avg	0.837	0.843	0.839	1500
Weighted avg	0.845	0.842	0.843	1500

H.4 AUC, FPR, FNR, and ROC

In addition to standard accuracy and F1 scores, we computed the area under the receiver operating characteristic curve (AUC), false positive rate (FPR), and false negative rate (FNR) for each classification model across platforms. The tables below summarize these measures for the out-of-sample validation data.

On the Twitter validation set, AUC values range from 0.94 to 0.99, indicating excellent discrimination. False positive rates are uniformly low (< 0.13), and false negative rates range from 0.11 to 0.32, suggesting that errors are generally conservative in nature—models occasionally miss relevant cases rather than over-predicting categories.

Performance on the Facebook validation set is nearly identical, with AUC values between 0.93 and 0.99 and similarly low FPRs and FNRs, confirming that the classifiers generalize effectively to another short-form, social-media context.

Newsletter classifications, which involve longer and more heterogeneous text, exhibit slightly lower but still strong performance, with AUC values ranging from 0.91 to 0.97. The modest increase in false negative rates for some categories reflects the greater linguistic and stylistic variation of long-form communication, consistent with expectations for cross-domain application of transformer-based models. Overall, the results demonstrate stable and high discrimination across all message types and platforms.

TABLE H.19: MODEL DISCRIMINATION METRICS (TWITTER VALIDATION SET)

Category	AUC	FPR	FNR
Claiming Combined	0.959	0.026	0.319
Service	0.981	0.032	0.115
Bipartisan	0.993	0.003	0.117
Negative Partisan	0.976	0.029	0.152
Position Taking	0.953	0.132	0.107
Policy Claiming	0.963	0.021	0.264
Advertising	0.937	0.094	0.220
Range:	0.937–0.993	0.003–0.132	0.107–0.319
Mean:	0.966	0.048	0.185

Note: Table displays the area under the ROC curve (AUC), false positive rate (FPR), and false negative rate (FNR) for each classification model evaluated on the manually coded Twitter validation set. High AUC values across all categories indicate strong discriminatory performance and low rates of misclassification.

TABLE H.20: MODEL DISCRIMINATION METRICS (FACEBOOK VALIDATION SET)

Category	AUC	FPR	FNR
Claiming Combined	0.957	0.035	0.275
Service	0.988	0.043	0.038
Bipartisan	0.981	0.006	0.204
Negative Partisan	0.968	0.035	0.143
Position Taking	0.941	0.128	0.117
Policy Claiming	0.961	0.024	0.273
Advertising	0.931	0.084	0.276
Range:	0.931–0.988	0.006–0.128	0.038–0.276
Mean:	0.961	0.051	0.189

Note: Table displays the area under the ROC curve (AUC), false positive rate (FPR), and false negative rate (FNR) for each classification model evaluated on the manually coded Facebook validation set. Model discrimination remains high across categories, comparable to Twitter.

TABLE H.21: MODEL DISCRIMINATION METRICS (NEWSLETTERS VALIDATION SET)

Category	AUC	FPR	FNR
Claiming Combined	0.948	0.020	0.342
Service	0.962	0.071	0.128
Bipartisan	0.963	0.003	0.190
Negative Partisan	0.974	0.028	0.214
Position Taking	0.924	0.165	0.149
Policy Claiming	0.954	0.014	0.434
Advertising	0.911	0.079	0.355
Range:	0.911–0.974	0.003–0.165	0.128–0.434
Mean:	0.948	0.054	0.259

Note: Table displays the area under the ROC curve (AUC), false positive rate (FPR), and false negative rate (FNR) for each classification model evaluated on the manually coded newsletters validation set. Although overall discrimination remains strong (AUCs above 0.91 for all categories), false negative rates are somewhat higher than for Twitter or Facebook, reflecting the greater heterogeneity and length of newsletter content. These patterns are consistent with expectations for cross-domain application of models trained primarily on social media text.

I Tables of Full Regression Estimates

TABLE I.1: ONLINE MESSAGING AND BEHAVIORS IN CONGRESS (ALL ESTIMATES)

	Credit Claiming		DV: Percent of Messages Classified as...				Constituent Service	
	(1)	(2)	Bipartisanship (3)	(4)	Negative Partisanship (5)	(6)	(7)	(8)
Bills Introduced	0.060*** (0.018)	0.068*** (0.017)						
Percent Opposite-Party Cosponsored Bills			0.108*** (0.011)	0.080*** (0.011)	-0.360*** (0.023)	-0.311*** (0.026)		
Local Mentions in Floor Speeches							1.088** (0.377)	0.777* (0.363)
Female		0.647 (0.512)		0.177 (0.156)		-0.961+ (0.570)		0.121 (0.755)
African American		-2.285*** (0.652)		-0.228 (0.195)		-3.427*** (0.883)		-0.989 (1.224)
Hispanic		-0.001 (0.920)		-0.381 (0.335)		-2.134** (0.764)		1.802 (1.310)
District Partisanship		-16.391*** (2.546)		-4.966*** (0.866)		14.723*** (3.184)		-15.994*** (3.847)
General Election Vote Share		-0.036* (0.015)		-0.014** (0.005)		-0.005 (0.019)		-0.054+ (0.028)
Party Leader		-2.230*** (0.587)		0.326 (0.281)		2.795*** (0.822)		-2.694** (0.940)
Committee Chair		-1.671* (0.713)		0.276 (0.274)		0.058 (0.899)		-1.421 (1.278)
Seniority		0.107* (0.052)		-0.004 (0.021)		0.146* (0.064)		-0.271*** (0.069)
Num. Obs.	2,099	2,054	2,099	2,054	2,099	2,054	1,503	1,479
Session	113-117	113-117	113-117	113-117	113-117	113-117	111-114	111-114
Controls	N	Y	N	Y	N	Y	N	Y
Party-Session FEs	Y	Y	Y	Y	Y	Y	Y	Y

Note: Estimates are from OLS regression models. Standard errors are clustered by member. * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

TABLE I.2: FULL TABLE OF ESTIMATES FOR FIGURE 6

	Likes	Twitter Retweets	Replies	Likes	Facebook Shares	Comments
Advertising	-0.188** (0.008)	-0.182** (0.008)	-0.059** (0.006)	-0.178** (0.008)	-0.289** (0.009)	-0.107** (0.006)
Constituent Service	-0.268** (0.008)	-0.205** (0.006)	-0.202** (0.006)	-0.332** (0.008)	-0.277** (0.008)	-0.429** (0.007)
Credit Claiming	-0.108** (0.009)	-0.142** (0.008)	-0.104** (0.007)	0.015* (0.007)	-0.125** (0.008)	-0.093** (0.006)
Position Taking	0.127** (0.008)	0.289** (0.008)	0.281** (0.006)	0.064** (0.008)	0.263** (0.009)	0.497** (0.008)
Negative Partisanship	0.471** (0.017)	0.588** (0.017)	0.580** (0.013)	0.188** (0.011)	0.623** (0.012)	0.657** (0.013)
Bipartisanship	-0.039** (0.009)	-0.038** (0.008)	-0.017* (0.008)	-0.078** (0.007)	-0.139** (0.007)	-0.027** (0.008)
Text Partisan Extremity	0.048** (0.018)	0.151** (0.017)	0.005 (0.014)	-0.003 (0.012)	0.102** (0.014)	0.066** (0.013)
Num.Obs.	4,712,478	4,712,478	4,712,478	2,427,335	2,427,335	2,427,335
R2 Adj.	0.718	0.607	0.649	0.552	0.513	0.603

Table displays coefficients from OLS models. Standard errors clustered by member shown in parentheses. *p<0.05; **p<0.01

J Within-Member Polarization Estimates

TABLE J.1: WITHIN-MEMBER POLARIZATION ESTIMATES

	DV: Partisan/Ideological Extremity					
	Rhetoric		Roll Call Votes		Contributions	
Session	-0.079** (0.025)	0.404*** (0.019)	0.001 (0.014)	-0.021 (0.011)	0.107*** (0.012)	-0.020 (0.014)
Majority Party Member	0.189* (0.076)	0.226*** (0.061)	0.029 (0.046)	0.014 (0.033)	-0.027 (0.034)	-0.097* (0.043)
Copartisan President	0.603*** (0.075)	0.156** (0.050)	-0.024 (0.034)	0.010 (0.023)	0.020 (0.018)	-0.170*** (0.025)
Num.Obs.	1,487	1,630	1,487	1,630	1,487	1,630
Party	Democrats	Republicans	Democrats	Republicans	Democrats	Republicans
Member FEs	Y	Y	Y	Y	Y	Y

Note: Estimates are from OLS regression models. Standard errors are clustered by member. * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

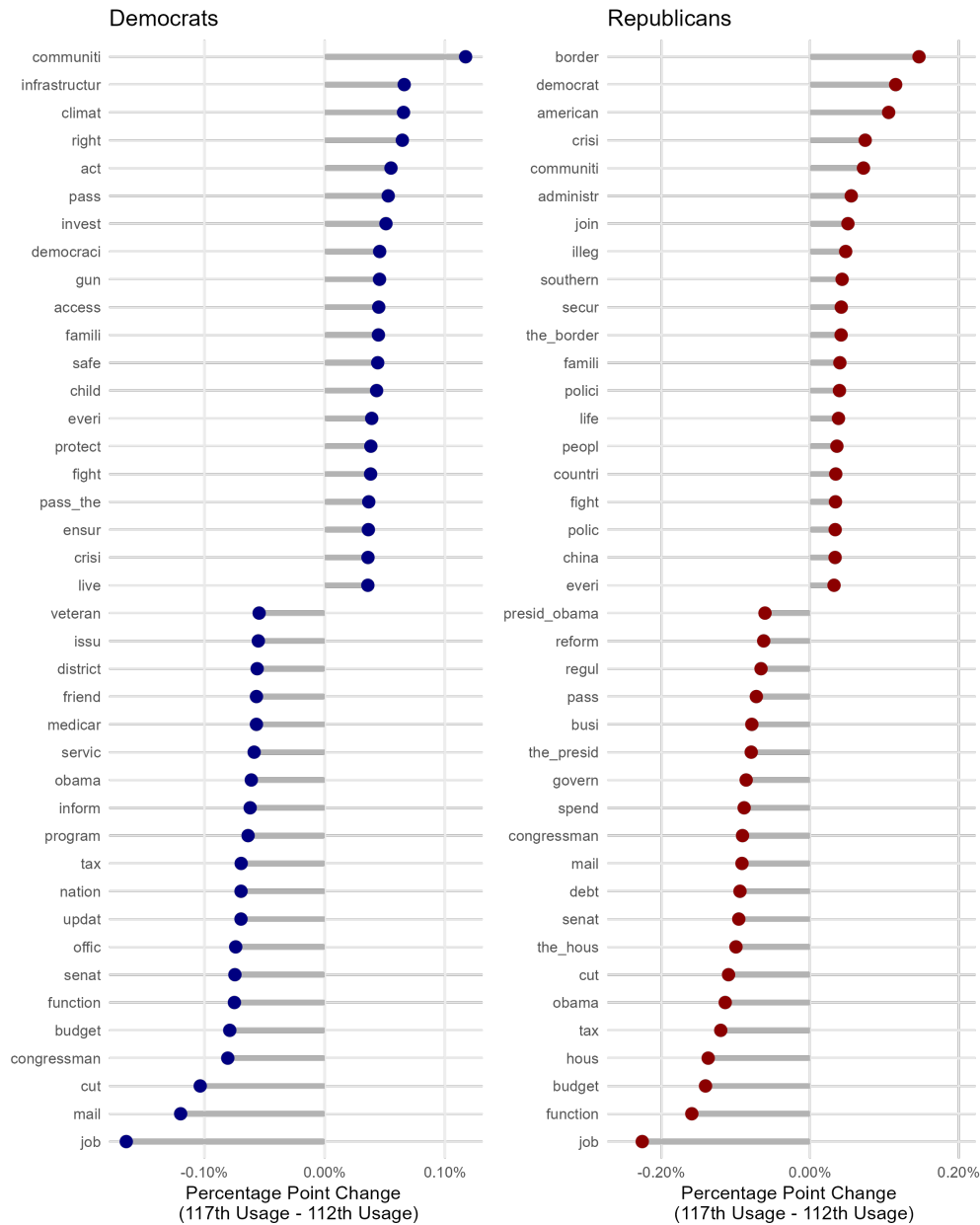
K Rhetorical Polarization Trends, Naive Bayes Measure

TABLE K.1: ESTIMATED POLARIZATION TRENDS USING NAIVE BAYES SCALING

	DV: Partisan Extremity (Naive Bayes)			
	Democrats		Republicans	
Session	0.064** (0.021)	0.028 (0.020)	0.144*** (0.021)	0.150*** (0.020)
Majority Party Member	-0.167* (0.068)	-0.224*** (0.059)	-0.105 (0.072)	-0.184** (0.063)
Copartisan President	0.685*** (0.065)	0.652*** (0.063)	0.138** (0.052)	0.185*** (0.054)
Num.Obs.	1,450	1,450	1,575	1,575
Member FEs.	N	Y	N	Y

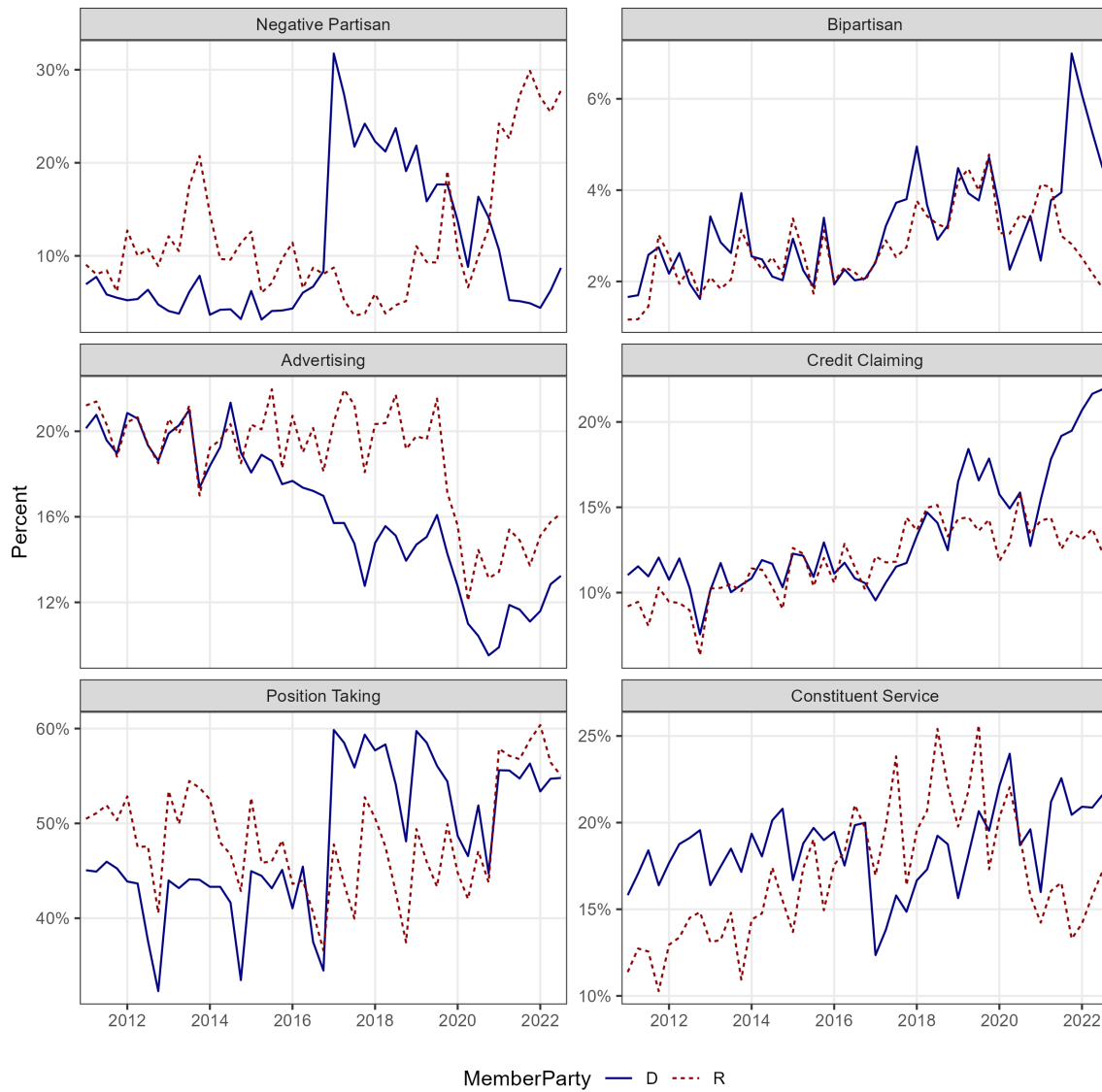
L Changes in Word Usage Across Time Period

FIGURE L.1: CHANGES IN WORD USAGE BY PARTY, 112TH TO 117TH SESSIONS



M Message Type Trends, by Party

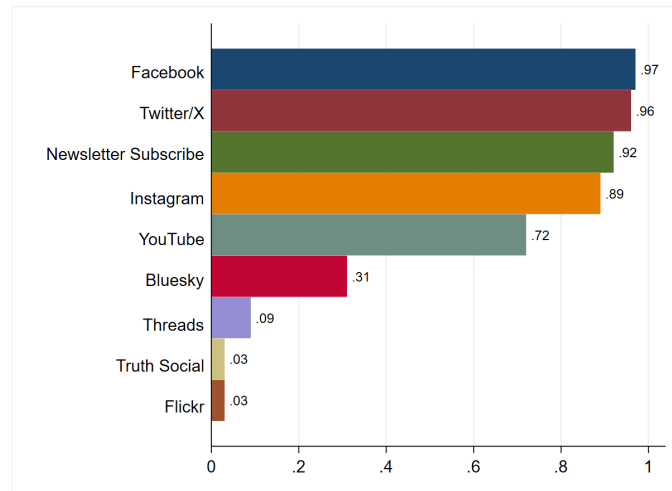
FIGURE M.1: TRENDS IN PARTISAN AND REPRESENTATIONAL CATEGORIES, BY PARTY (2011 - 2022)



Note: The figure displays trends in the quarterly averages of the two partisan and four representational categories trends described above, displayed separately for each of the two major parties between 2011 and 2022.

N Social Media and Newsletter Links on House Members' Homepages

FIGURE N.1: PERCENTAGES OF HOUSE MEMBER HOMEPAGES WITH LINKS TO EACH TYPE OF MEDIA



Note: Percentages calculated based on the authors' August 2025 analyses of all U.S. House members' official homepages. Platforms linked by <3% of members are omitted.