

IMAP: A Tweet-Based Index of Messaging and Affective Polarization among Partisan Elites

Sam Frederick*

This version: September 29, 2025[†]

Abstract

Recent research has advanced our understanding of mass-level affective polarization, defined in terms of partisans' identity-based feelings toward the two major parties in the U.S. Mass-level affective polarization has risen in recent decades and appears to carry real consequences for opinions and behavior. Less is understood, however, about whether political elites are similarly affectively polarized and the consequences this carries for elite behavior and representation. Recent events, including the January 6th Capitol Insurrection, as well as scholarship on the behavioral consequences of partisanship for the public, have laid bare the need to understand partisan animosities among politicians. In this paper, I propose a novel framework for measuring elite partisan messaging and affective polarization which I call the Index of Messaging and Affective Polarization (IMAP). IMAP captures both the prominence and direction of sentiment toward the two major parties in tweets posted by elected officials, combining machine learning methods with Bayesian item-response models. After describing the measure, I demonstrate that the measure comports with the reputations of prominent national political figures as extreme partisans. Next, I illustrate how IMAP relates to a variety of behavioral outcomes, including congressional bill co-sponsorship decisions, public responses to partisan scandals, and mass-level affect. IMAP also appears to capture the behavior of American governors as well. Finally, I conclude by considering potential avenues for future research employing IMAP.

*Ph.D. Candidate, Department of Political Science, Columbia University; sdf2128@columbia.edu

[†]This draft is a work in progress. Please do not circulate without the author's permission.

Introduction

On January 6, 2021, supporters of former President Trump broke into the United States Capitol. Inside, they attempted to halt the ceremonial counting of the electoral college votes from the 2020 presidential election, vandalizing the building and assaulting members of the Capitol Police. Following the November election, Trump had been feeding a steady stream of misinformation about the election to his followers. He claimed repeatedly that the election was stolen from him by Democrats. On the day of the riot, Trump held a rally in D.C., urging his supporters to descend on the Capitol, and his lawyer, Rudy Giuliani, encouraged them to engage in “trial by combat.” Ultimately, the stream of misinformation and provocations culminated in an outbreak of political violence and an attack on the United States Capitol not seen in modern history. The Capitol insurrection was not the only example in recent years of politicians’ words having repercussions for mass political behavior. At campaign rallies, Trump often openly encouraged his supporters to “knock the hell out of” protesters—and protesters were beaten on multiple occasions ([Cineas 2021](#)).

More recently, some Democrats in Congress have reported feeling afraid of their counterparts in the Republican Party, and metal detectors were installed outside of the House chamber to prevent members of Congress (hereafter MCs) from carrying guns onto the floor—though several GOP politicians have skirted the detectors ([Sarlin 2021](#)). Republican Representative Paul Gosar was censured by the House and stripped of his committee memberships after tweeting a video which showed him, in anime form, killing Democratic Representative Alexandria Ocasio-Cortez and threatening President Joe Biden ([Quinn 2021](#)). Additionally, several Republicans have been subjected to a barrage of criticism from their co-partisans in Congress, including an accusation that they are “traitors” by Republican Representative Marjorie Taylor Greene, for simply voting with their partisan opponents. Their offices have received death threats from the public as well following the vote ([Edmondson 2021](#)).

Undoubtedly part of this trend toward violence, scholars have recently documented a striking increase in affective polarization at the mass level. Affective polarization is defined

by dislike of the opposing party and warm feelings toward one’s own party among those who identify as partisans. Partisans in the mass public are more likely to say they would be upset if a family member married someone from the other party (Iyengar, Sood and Lelkes 2012), and they are even willing to discriminate against opposing partisans in lab experiments (Iyengar and Westwood 2015) and on the job market (Gift and Gift 2015). Further, Iyengar and Westwood (2015) show that this bias has even seeped into the mass public’s implicit attitudes. From the mass behavioral literature, then, it appears that affect toward the parties has consequences for behavior and that this concept is distinct from ideology (e.g., Iyengar, Sood and Lelkes 2012).

Despite increasing evidence of affective polarization driving mass behavior, there has been less work examining the affective polarization of elites. With no consistent measure of elite affective polarization, we have no reliable way of discerning the influence of affective polarization on behavior at the elite level, and it is challenging to understand the dynamics of affective polarization between elites and masses. The intensity of former President Trump’s exhortations to violence against and denigration of political opponents highlights particularly well the need to understand affective polarization among elites.

It is toward this end that I develop and validate a framework for measuring affective polarization in political discourse called the Index of Messaging and Affective Polarization (IMAP). Given the recent partisan ire expressed on Twitter, most prominently by former President Trump, I employ the tweets of prominent U.S. politicians for development and validation of the framework—though, in principle, this framework could be applied to other forms of communication like floor speeches and press releases. I propose examining the intensity of partisan sentiment expressed by politicians on Twitter as a means of indexing elite affective polarization as well as the closely related concept of partisan messaging. My framework of measurement utilizes natural language processing to categorize a large number of texts from politicians and deploys Bayesian Item Response Theory to capture the overall partisan affect of each politician. The measures developed here are significantly related

to various forms of elite behavior, even after accounting for existing measures of ideology. Specifically, elites who are more devoted to affectively polarized party messaging appear to be less bipartisan and more frequent features on Fox News. Among Republicans, affective polarization in messaging is significantly related to Trump loyalty and perceptions of conservatism by Republican activists.

In the next section, I describe IMAP and justify the underlying assumptions required to accept this measure. I also demonstrate the face validity of this measure. Then, I validate this measure by comparing it to various forms of congressional behavior which are plausibly influenced by partisan affect. Finally, I conclude with potential future applications of my measure to the study of politics.

Index of Messaging and Affective Polarization

Previous Literature

My measure of partisan messaging and affective polarization spans two established literatures: institutional scholarship on partisan messaging and behavioral work examining affective polarization. Using tweets posted by U.S. politicians, I capture whether these tweets mention either or both parties and the direction of the partisan messages within partisan tweets. From these partisan tweets, I derive the overall partisan orientation of each politician’s Twitter feed. I use this orientation on Twitter as a behavioral measure of affect toward the parties and of commitment to party messaging.

A flourishing literature on mass-level polarization has demonstrated that, while members of the American public are not broadly ideologically polarized, they are nonetheless *affectively* polarized (Iyengar, Sood and Lelkes 2012; Levendusky 2009). Partisans increasingly hate their opponents and are willing to discriminate against them in a variety of settings (Gift and Gift 2015; Iyengar and Westwood 2015; McConnell et al. 2018). That members of the mass public detest their partisan opponents and that this carries consequences for behavior have been largely absent from literature on elites—in part, due to the lack of a systematic

measure of elite affect. Still, the potential for real behavioral consequences, demonstrated at the mass level, indicate that this area should not be neglected as elite affective polarization could have important implications for representation and democratic governance. For example, affectively polarized elites, like their citizen counterparts, might discriminate against opposing partisans in government and among their constituents, meaning that cross-partisan lawmaking and representation might suffer. Before questions like this can be answered, it is crucial to have a systematic measure of elite affect. In this paper, I lay out such a measure which can be used to examine the consequences of affect for representation and governance.

Beginning with the work of [Cox and McCubbins \(1993, 2005\)](#), scholars have noted the importance of party “brands” to behavior in Congress. Political parties are incentivized to distinguish themselves from their partisan opponents and to cultivate a favorable image for their own party through party messaging. In particular, [Lee \(2016\)](#) notes the centrality of increasing competition for control of Congress to the growing prevalence of partisan conflict in messaging. Moreover, [Iyengar, Sood and Lelkes \(2012\)](#) hypothesize that this partisan messaging is an important driver of mass-level affective polarization, and [Mutz and Reeves \(2005\)](#) show that incivility in elite discourse causes decreases in trust in government. Thus, the partisan messages sent by elites have potentially great consequences for mass-behavior and democratic functioning. Given limitations on the number of texts it was possible to code for partisan messaging, however, previous analyses were limited in their scope to small samples of communications from small groups of politicians (e.g., [Lee 2016](#)). Recent advances in machine learning and text analysis have opened the door the analysis of larger corpuses of texts from a variety of political actors. I employ these methods to analyze the Twitter feeds of members of Congress for the presence of partisan messages, making it possible to more rigorously study elite partisan messaging. While I believe this tweet-based measure of affect in partisan messaging can prove valuable to the studies of partisan messaging and affective polarization, such a behavioral measure relies on a crucial justification which I discuss in the next section.

Measuring Affective Polarization Using Twitter

To reliably measure prioritization of partisan messaging and the affective orientations of officeholders from tweets, I need to assume that partisan tweet patterns reflect the true underlying affect of their posters toward the parties.

A long line of political science research, dating back to [Fenno \(1977, 1978\)](#), suggests that MCs tend to develop unique “home style[s]” which they use to communicate with their constituencies. Importantly, a member’s presentation of self is largely determined by her personality. Further, while party leadership in Congress can often control the floor agenda and at times, influence member voting ([Cox and McCubbins 2005](#); [Rohde 1991](#)), they have relatively fewer incentives to control member messaging. In fact, party leadership has an incentive to allow MCs to develop individualized patterns of communications (unique home styles) to aid member reelection, as individual reelections help the leadership’s broader goal of maintaining power for the party. Thus, while votes against the party can carry significant policy repercussions and prominent party defeats on legislation caused by internal dissension can weaken the party brand ([Cox and McCubbins 2005](#)), individual member tweets are relatively low cost. Indeed, according to [Lee \(2016\)](#), parties often gauge the success of messaging strategies by *whether* their members adopt the message, reflecting member discretion in messaging. It therefore seems reasonable to assume that politicians’ tweets are relatively free from the constraints of party leadership.

That said, MCs are likely still subject to the constraints of constituency. Marginal politicians who personally hate the opposing party might conceal that hatred in their public comments given their reliance on voters from the other party for support. This is a potential problem without an obvious solution; however, I would note that members from different parties representing politically similar constituencies have vastly divergent voting records, so, while there is some moderation in voting records based on constituency, it appears that constituencies are not wholly determinative of voting ([Ansolabehere, Snyder and Stewart 2001](#); [Bafumi and Herron 2010](#)). If voting leaves room for the influences of party leadership,

ideology, and other idiosyncratic factors, it seems likely that the patterns of communication in tweeting would also not be fully determined by constituency.¹ Further, given the necessity of making cross-party appeals to voters in marginal districts, politicians who truly abhor the other party and its members would logically have a disincentive to run in marginal districts, perhaps selecting out of representing constituencies that would force them to appeal to the other party’s voters and to govern cooperatively with the other party. In short, while strategic considerations certainly matter for tweets, following the findings of [Fenno \(1978\)](#), presumably politicians make those strategic considerations *conditional on their “type”*—in this case, conditional on their affect. Politicians who are not extremist partisans but who adopt extreme partisan self-presentations would likely experience cognitive dissonance and should seek to resolve that dissonance.

At a minimum, even if the above assumption does not hold, the measure I develop would still stand to contribute to our understanding of partisan messaging.

Overview of IMAP

Data

I collected tweets from various U.S. politicians with Twitter’s Academic API which allows users to collect more tweets from Twitter’s archive than previously available. To create a manageable and defined group of accounts, I gathered only the official Twitter account handles of every MC who served during the 116th Congress (accounts that are either linked to from MCs’ .gov websites or contain a link to the .gov website in the bio). At the time of collection, some accounts of retiring or defeated members had already been deleted, generating some missingness in the data.² I also compiled a list of the official accounts of sitting U.S. governors. Using these two lists of Twitter handles, I downloaded all of the tweets from each

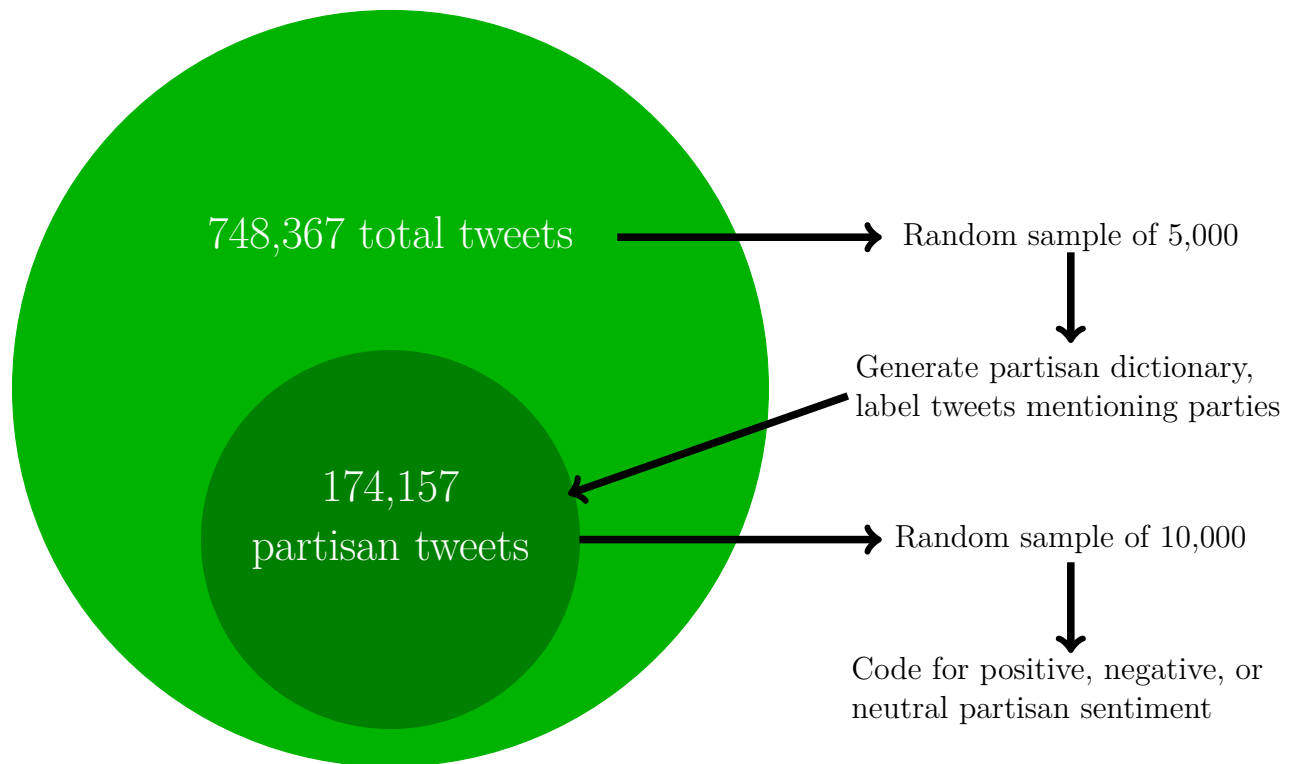
¹Indeed, I show in [Appendix D](#) that my measures of partisan affect have low correlations with traditional measures of ideology and constituency partisanship.

²Upon retirement or losing reelection bids, MCs often delete their official accounts. Because many members of Congress retired or lost their reelection bids in 2018, the comprehensiveness of the data would be greatly reduced by extending it further back in time.

user’s timeline posted between January 3, 2019 and January 3, 2021 (the beginning and end of the 116th Congress). I excluded non-English tweets as translation apps can be quite inaccurate. Importantly, there were very few non-English tweets, and most appeared to repeat the content of English tweets in different languages. I excluded retweets because, while the decision to retweet a message certainly conveys information, the information conveyed is not the politician’s original thought, and it is unclear whether the decision to retweet a message indicates agreement, disagreement, or something else. Indeed, many politicians explicitly note in their Twitter bios that retweets are not equivalent to endorsements of the views in the tweet. Next, I arranged the tweets by date and joined the tweets by their conversation ID, which indicates that the tweets are a part of the same text, or thread. I joined individual tweets belonging to each thread into a single text because the purpose of threads is to discuss one cohesive topic (i.e., threads are intended to be read in whole rather than in part). The vast majority of tweets in the dataset were only single tweets (89.54 percent) rather than threads. The maximum number of tweets in a single thread was 132. In total, I began with 748,367 tweets from 554 users, each of whom sent about 1,351 tweets on average.

Having generated a dataset of complete tweet threads, I read through a random sample of 5000 tweets to compile dictionaries of the terms most frequently used to describe each party (i.e., the parties themselves, party leadership, and the parties’ presidential candidates or incumbent). I combined these dictionaries with regular expressions to label tweets mentioning the Democratic Party and/or the Republican Party and examined only this subset of tweets mentioning one or both of the parties. Having labelled partisan tweets, I drew a random sample of 10,000 from the subset of partisan tweets. To develop a measure of affective polarization and partisan messaging, I hand-coded these 10,000 tweets by hand for the presence or absence of negative and positive affect toward the referenced party. The dictionaries used to generate this sample of partisan tweets was somewhat imperfect, so I updated the dictionaries as well to better capture partisan references. The full regular-expression dictionaries can be found in [Appendix A](#). In total, after updating the dictionaries, I had

Figure 1. Hand-Coding Process



hand-coded a sample of 9,823 tweets from a total of more than 174,000 partisan tweets. The full hand-coding process is depicted in [Figure 1](#).

Tweets that were categorized as negative blamed the referenced party for some potential or realized negative outcome such as inaction, corruption, loss of health insurance, or a government shutdown (often invoking "valence issues," as described by [Stokes 1963](#)). Negative partisan tweets also attacked the referenced party for its issue positions. Tweets that were categorized as positive toward a party fit within [Mayhew's \(1974\)](#) categories of advertising, credit-claiming, and position-taking—but applied to the parties rather than to individual members. Thus, positive tweets claimed credit on behalf of a party (e.g., Republicans claiming credit for a good economy following the passage of the Tax Cuts and Jobs Act), advertised the party's good valence attributes (e.g., Democrats holding a hearing to determine the best way to address rising health insurance costs), or emphasized the party's

Table 1. Distribution of Hand-Coded Tweets

	Democratic Tweets			
	Negative	Positive	Neutral/Nonpartisan	Total
Number	1363	1478	261	3102
Percent of Democratic Tweets	43.94	47.65	8.41	
Percent of Total Tweets	13.88	15.05	2.66	
	Republican Tweets			
	Negative	Positive	Neutral/Nonpartisan	Total
Number	5095	2083	834	8012
Percent of Republican Tweets	63.59	26	10.41	
Percent of Total Tweets	51.87	21.21	8.49	
	Tweets Mentioning Both Parties			
	Negative Democratic	Positive Democratic	Neutral/Nonpartisan Democratic	
Negative Republican	11	509	71	
Positive Republican	499	38	15	
Neutral/Nonpartisan Republican	41	39	68	

position on an issue favorably (e.g., Democrats touting support for criminal justice reform following George Floyd’s murder or highlighting the potential benefits of a proposed policy). This coding framework is in line with that used by [Lee \(2016\)](#) to analyze press releases from party leadership, applying [Mayhew’s \(1974\)](#) categories to partisan communications. Neutral or nonpartisan tweets were the residual category, capturing tweets which had no clear directionality, were nonpartisan, or did not reference the parties. The categories were mutually exclusive, meaning that a tweet could not be coded as both negative and positive toward the same party—though a single tweet could be positive toward one party and negative toward the other. In [Table 1](#) below, the summary results of manual coding of the sample of 9,823 tweets are displayed. Few tweets fall into the neutral or nonpartisan category—only about 10 percent of both Democratic and Republican tweets. There were far more tweets referring to the Republican Party than to the Democratic Party, and in fact, the majority of all tweets were negative tweets directed at the Republican Party. Additionally, in tweets mentioning both parties, most tweets were clearly positive toward one party and negative toward the other: there were few tweets mentioning both parties that expressed clear sentiment toward one party and were neutral toward the other or that expressed the same sentiment toward both parties.

Methods

Given that my dataset contains more than 174,000 tweets referencing one or both of the parties, I apply natural language processing methods to categorize the remaining tweets which allows me to analyze far more tweets than would be possible by hand. Following [Grimmer and Stewart \(2013\)](#), I preprocessed the tweets by removing punctuation and numbers—though due to the nature of my data, I left hashtags and at-symbols in the text. I removed hyperlinks from the text. Then, I removed stopwords, or words that appear commonly in text and carry little discriminating meaning on their own (like “the” and “is”). Finally, I stemmed the words using the Porter Stemmer which is meant to reduce similar words to a common root form, so that, for example, plural versions and singular versions of the same words are treated as the same word.

Machine learning algorithms require that text be represented in vector form. There are several ways of transforming texts to vectors. To maximize predictive accuracy of the machine learning algorithms, I tested several different methods of transforming the preprocessed tweets to vector representations, in all cases omitting corpus-specific stopwords (those which occur in more than 99 percent of all texts in the corpus). I tried standard Bag of Words representations with unigrams which transform texts to vectors by counting the frequency of each individual word’s occurrence in a text. However, Bag of Words can lose the context of words, so I also transformed the tweets into bigrams (each unique two-word sequence) and unigrams plus bigrams simultaneously to test whether predictive accuracy of the models was improved by the additional context. In addition to simple Bag of Words representations, I used Term Frequency-Inverse Document Frequency (TFIDF) weighting which weights words by the inverse of their prevalence across documents, such that more frequent words are given lower weights and less frequent words are assigned higher weights. Finally, to more fully capture the context of a document, I employed Google’s Doc2Vec, which transforms texts to vectors (in my case, of size 300) by attempting to predict random words from a text and adjusting document vector weights to best accomplish this task ([Le and Mikolov 2014](#)). I

computed the document vectors using Gensim, a Python library from [Řehůřek and Sojka \(2010\)](#). In total, I fit models using seven different text vectorizations: BOW unigrams, BOW bigrams, BOW bigrams and unigrams, TFIDF unigrams, TFIDF bigrams, TFIDF bigrams and unigrams, and Doc2Vec.

After transforming tweets to vectors, I proceeded to fit a variety of machine learning algorithms to automatically perform the partisan sentiment categorization. I fit separate models on the tweets which mentioned Democrats and those which mentioned Republicans. I employed an 80-20 train-test split to evaluate the performance of the machine learning algorithms (I trained the algorithms on 80 percent of the hand-coded tweets mentioning each party and tested model performance using the remaining 20 percent of the hand-coded tweets). [Grimmer and Stewart \(2013\)](#) note that this is the best way to validate the performance of classification models: it allows researchers to see how the model performs on texts unseen by the model but which have been classified by the researcher to gain an understanding of how the classifier might work on the unseen and unclassified texts.

I chose models which have been widely applied in machine learning and text classification with great success. Then, I tested model performance using the 20-percent test data samples. I used a variety of metrics to evaluate model performance on the test data: simple accuracy, which evaluates the proportion of classifications from the model which were correct; precision for each category, which evaluates the extent to which predicted categorizations reflect actual categorizations; recall for each category, which evaluates the fraction of true categorizations predicted by the model (i.e., whether the model misses texts which are actually in a category); F1-score, which is the harmonic mean of precision and recall for each category ([Pedregosa et al. 2011](#)); and finally, balanced accuracy, which averages the recalls for each category, providing a clearer picture of recall when classes have different sizes, as they do in the coded sample of tweets.

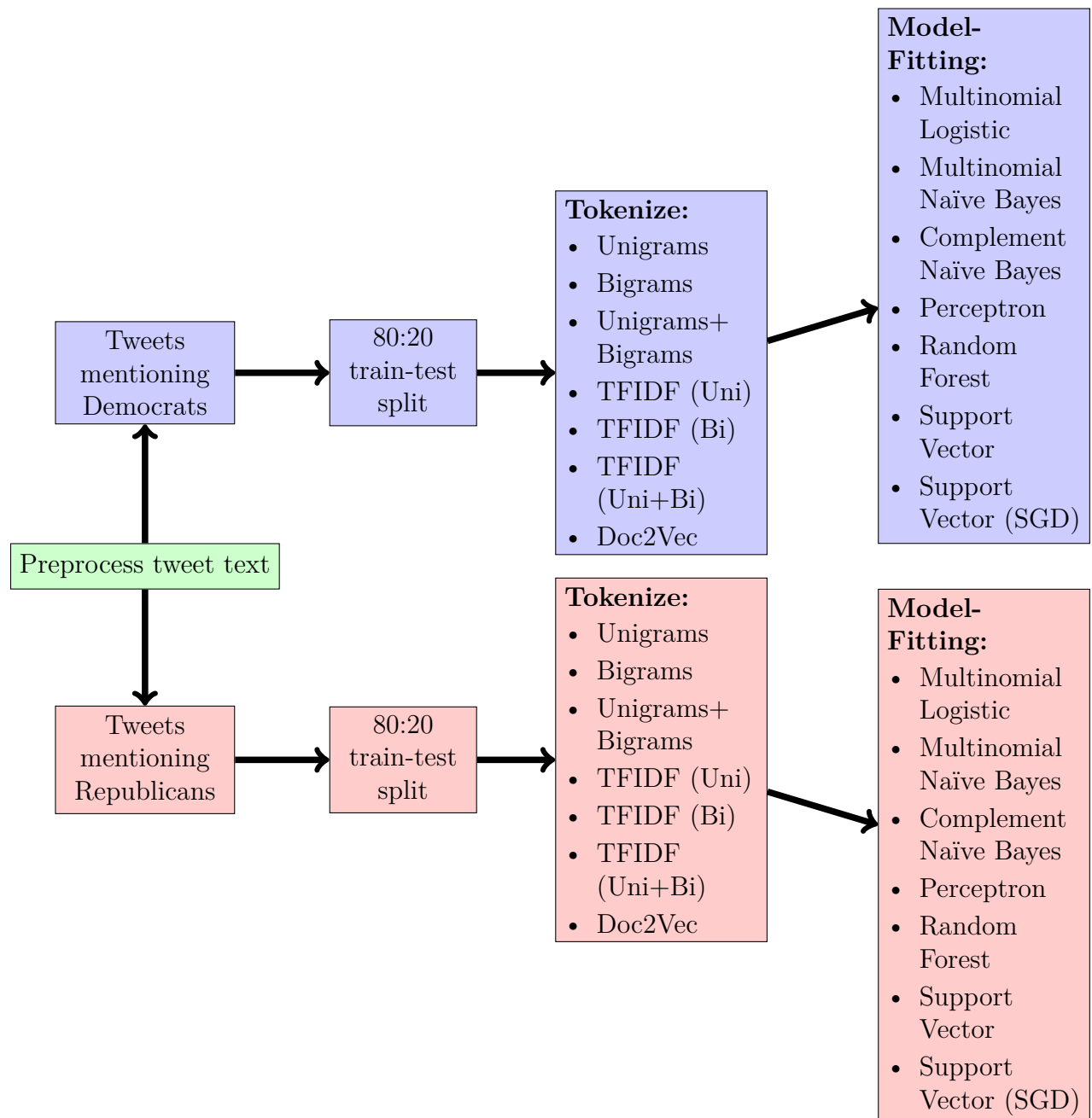
First, I fit multinomial logistic regression models using each of the text vector forms mentioned above to predict text categorizations. Generally, the multinomial regression mod-

els performed quite well—even against more sophisticated machine learning models. Next, I used a variety of Naive Bayes classifiers, often employed in text analysis. Naive Bayes applies Bayes’ Rule to the text vectors and categories ([Jurafsky and Martin 2021](#)). In essence, it assumes that the probability of a text belonging to a category given its word vector can be represented by the probability of its being in that category multiplied by the probability of that word vector conditional on the category. It is "naive" in that it unrealistically assumes words within a document are conditionally independent, given the category. I employed both Multinomial Naive Bayes and a variant classifier, Complement Naive Bayes. Complement Naive Bayes is a type of Naive Bayes classifier developed by [Rennie et al. \(2003\)](#), which is intended to perform better than Multinomial Naive Bayes on imbalanced datasets (datasets with unequal numbers of documents in each class, like the hand-coded data). Rather than computing the probability of a word given a class, Complement Naive Bayes takes the probability that a word appeared in every class *but* the estimated class (the complement) and aims to minimize the complement probability.

Additionally, I fit random forest classifiers which draw randomly with replacement from the data, samples random subsets of the text vectors, and trains decision tree models using the random samples of data and vector features ([Breiman 2001](#)). Ultimately, Random Forests vote for the proper category for a text. Next, I tested the Perceptron algorithm which assigns and updates weights to words in the vector to predict text class. Finally, I fit two versions of Support Vector Classifiers which attempt to maximize the distance between texts in categories and a dividing plane between categories ([Pedregosa et al. 2011](#)).

In all, I fit seven models using the seven text vectorizations, leaving me, ultimately, with total of 47 models fits each for tweets that mentioned Republicans and for those that mentioned Democrats, with my steps displayed in [Figure 2](#). Below, I display the results of the top-10 model fits, arranged in order of balanced accuracy scores, with the maximum for each metric in bold. Results for models fit on tweets mentioning Democrats are shown in [Table 2](#), and those fit on tweets mentioning Republicans are shown in [Table 3](#). The

Figure 2. Machine Learning Pipeline



models fit using bigrams alone performed worse than those fit with simple unigrams in almost every model and by almost every metric, so I omit bigrams test-set statistics in the tables below and in the complete test-set validation metrics in [Appendix B](#) for simplicity. Model predictions for the neutral/nonpartisan categories of tweets were much less accurate than for the positive and negative tweets. This could be due to the residual nature of the category which left tweets with no clear partisan directionality. It could also be due to the relatively small number of tweets which fell into the neutral/nonpartisan categories, which, in effect, left very few actual tweets in these categories in the test set (being incorrect by only a couple of tweets in the neutral/nonpartisan categories could have large impacts on precision and recall). For Republican tweets, Support Vector Machine trained with stochastic gradient descent with TFIDF-weighted unigrams and bigrams performed the best across a variety of metrics, including accuracy and balanced accuracy. Among the models for Democratic tweets, this Support Vector model with TFIDF-weighted unigrams and bigrams also performed quite well—only slightly below the highest balanced accuracy and above Perceptron’s overall accuracy. Indeed, on a variety of metrics the Support Vector Classifier performed the best of any model. For this reason, I opted to use Support Vector Classifiers with TFIDF-weighted unigrams and bigrams to categorize the remaining 164,384 uncategorized partisan tweets.

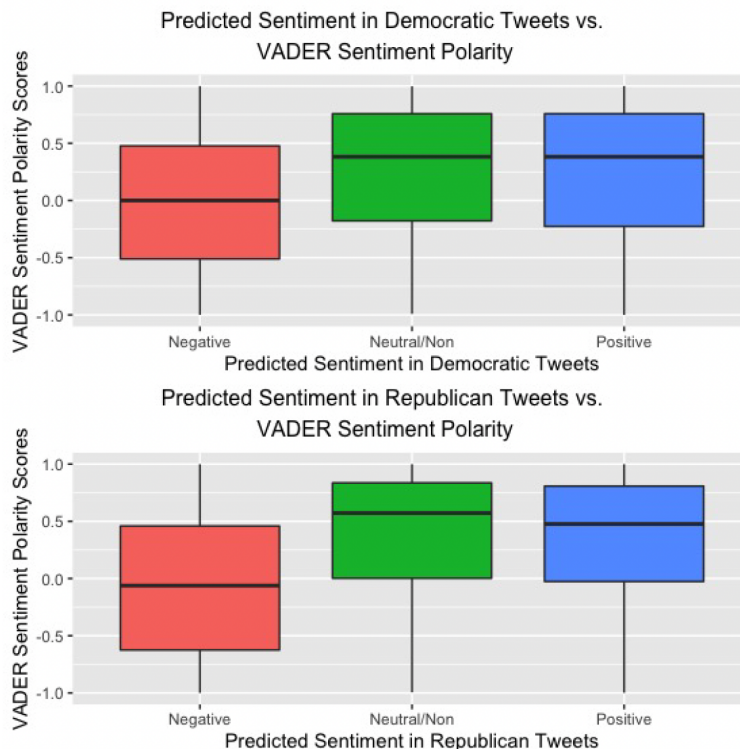
Table 2. Models for Tweets Mentioning Democrats

Vectorization	Model	Accuracy	Balanced Accuracy	Prec. Neg.	Prec. Neut.	Prec. Pos.	Rec. Neg.	Rec. Neut.	Rec. Pos.	F1 Neg.	F1 Neut.	F1 Pos.
TFIDF	SVC(SGD)	0.821	0.659	0.821	0.458	0.853	0.893	0.25	0.835	0.855	0.324	0.844
BOW	SVC(SGD)	0.805	0.661	0.831	0.406	0.822	0.859	0.295	0.828	0.845	0.342	0.825
TFIDF	Perceptron	0.824	0.662	0.832	0.55	0.835	0.883	0.25	0.853	0.857	0.344	0.844
Uni+Bi-BOW	SVC(SGD)	0.833	0.674	0.861	0.522	0.827	0.876	0.273	0.875	0.869	0.358	0.85
Uni+Bi-BOW	Perceptron	0.818	0.676	0.827	0.583	0.828	0.883	0.318	0.828	0.854	0.412	0.828
Doc2Vec	SVC(SGD)	0.792	0.677	0.834	0.333	0.832	0.846	0.386	0.799	0.84	0.358	0.815
Uni+Bi-BOW	Multinomial	0.829	0.678	0.853	0.52	0.83	0.896	0.295	0.842	0.874	0.377	0.836
Uni+Bi-TFIDF	SVC(SGD)	0.839	0.685	0.845	0.5	0.864	0.913	0.295	0.846	0.877	0.371	0.855
BOW	Multinomial	0.831	0.685	0.849	0.452	0.853	0.903	0.318	0.835	0.875	0.373	0.844
BOW	Perceptron	0.813	0.686	0.841	0.471	0.824	0.872	0.364	0.821	0.857	0.41	0.822

Table 3. Models for Tweets Mentioning Republicans

Vectorization	Model	Accuracy	Balanced Accuracy	Prec. Neg.	Prec. Neut.	Prec. Pos.	Rec. Neg.	Rec. Neut.	Rec. Pos.	F1 Neg.	F1 Neut.	F1 Pos.
Doc2Vec	Multinomial	0.792	0.66	0.856	0.55	0.694	0.904	0.392	0.683	0.879	0.458	0.689
Uni+Bi-TFIDF	Perceptron	0.82	0.666	0.92	0.85	0.642	0.895	0.188	0.915	0.907	0.308	0.754
BOW	Perceptron	0.795	0.672	0.888	0.451	0.696	0.895	0.409	0.714	0.891	0.429	0.705
BOW	ComplementNB	0.822	0.675	0.88	0.652	0.712	0.928	0.32	0.776	0.903	0.43	0.743
Uni+Bi-BOW	Perceptron	0.823	0.675	0.876	0.612	0.731	0.935	0.331	0.759	0.905	0.43	0.745
BOW	Multinomial	0.82	0.677	0.887	0.559	0.719	0.926	0.343	0.764	0.906	0.425	0.741
TFIDF	Perceptron	0.81	0.679	0.883	0.544	0.707	0.91	0.376	0.751	0.897	0.444	0.728
Uni+Bi-BOW	Multinomial	0.838	0.683	0.881	0.747	0.743	0.952	0.309	0.786	0.915	0.438	0.764
Uni+Bi-BOW	SVC(SGD)	0.835	0.685	0.88	0.667	0.75	0.946	0.331	0.776	0.912	0.443	0.763
Uni+Bi-TFIDF	SVC(SGD)	0.848	0.697	0.886	0.803	0.759	0.952	0.315	0.824	0.918	0.452	0.79

Figure 3. Sentiment Analysis of Tweet Categories



To ensure that the sentiment categorizations of my models are truly capturing underlying sentiments of the tweets, I performed a sentiment analysis of predicted partisan tweets from the models. [Hutto and Gilbert \(2014\)](#) compile a sentiment lexicon and combine it with language rules like negation, punctuation, and capitalizations to capture sentiment in tweets. Thus, by utilizing their lexicon and rules, implemented using their Python package, I retrieved their estimates of the positive, negative, or neutral sentiment directions in each partisan tweet. [Figure 3](#) displays a box plot of sentiment grouped by tweets’ predicted category for both Republican and Democratic tweets. Positive tweets—for both Republican and Democratic tweets—have higher sentiment scores than do negative tweets, as calculated using VADER. This indicates that predicted negative tweets are indeed more negative in their sentiments than predicted positive tweets. In [Appendix C](#), I further validate the model predictions with random samples of 10 tweets from each category.

Bayesian Item Response Theory

The predicted categorizations of the remaining partisan tweets from the trained algorithms allowed me to calculate the number of positive and negative Democratic tweets and the number of positive and negative Republican tweets posted each day by each user. Relying on the assumption, discussed above, that partisan tweeting patterns reflect the affect of the account holder toward the parties, I used Bayesian Item Response Theory to estimate two latent variables underlying the production of partisan tweets. Specifically, I modeled the probability that a user posted a negative or positive Democratic tweet as a function of their underlying Democratic affect plus any day-specific events, as certain days might generate higher or lower volumes of partisan tweets without reflecting the underlying preferences of the users. For example, when Congress is not in session, MCs might post less frequently—not because they are not polarized during this time but because they are not working in Congress at the time. I also modeled the probability that a user posted a negative or positive Republican tweet as a function of underlying Republican affect plus any day-specific events.

Thus, I recovered the latent variable, Democratic affect (which I refer to as a user’s Democratic Affect Score), using the following model:

$$Pr(DemTweet_{idt} = 1 | \delta_i, \gamma_d, DemDirection_{idt}) = \text{logit}^{-1}(\delta_i * DemDirection_{idt} + \gamma_d) \quad (1)$$

In this model, the probability that a user i posts a tweet mentioning the Democratic Party on day d is a function of their Democratic Affect Score, δ_i , the direction of the tweet, $DemDirection_{idt}$, and the day-specific events, γ_d . A negative Democratic Affect Score indicates that a user is more likely to post a negative Democratic tweet and less likely to post a positive Democratic tweet. To clarify, there are two outcome measures for each user on each day, for a total of 811,056 rows. Each day includes separate outcomes, indicating whether a user posted a positive Democratic tweet and whether the user posted a negative Democratic tweet. This model specification allowed users to post neither, either, or both positive and

negative partisan tweets. The probability of each outcome was described as a function of partisan affect and time shocks. For Republican tweets, I fit an analogous model:

$$Pr(RepTweet_{idt} = 1 | \theta_i, \pi_d, RepDirection_{idt}) = \text{logit}^{-1}(\theta_i * RepDirection_{idt} + \pi_d) \quad (2)$$

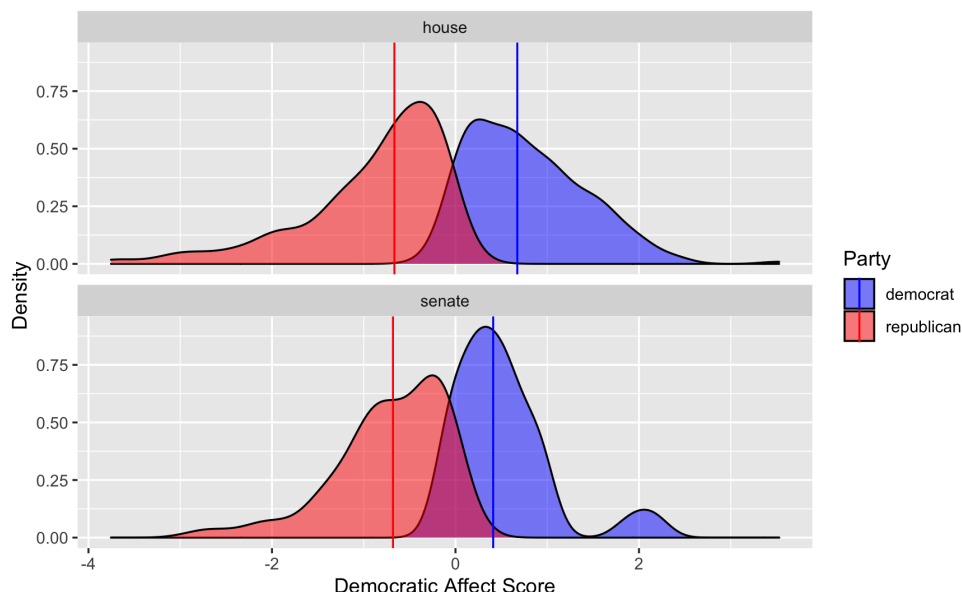
Here, θ_i is the Republican Affect Score, and π_d allows for date-specific events to influence the frequency of posts about Republicans. Finally, *RepDirection* is a binary variable that is either 1 or -1. Thus, politicians with positive Republican Affect Scores are more likely to post a positive Republican tweet on any given day and are less likely to post a negative Republican tweet.

Both θ_i and δ_i were given standard normal priors. I ran the model using Stan in RStudio. As [Clinton, Jackman and Rivers \(2004\)](#) and [Gelman and Hill \(2007\)](#) point out, standard item response theory models are not identified due to additive and multiplicative invariance. My approach ensures that those who post more often in a positive manner about Democrats will have positive Democratic Affect Scores and those who post more often in a negative manner about Democrats will have negative Democratic Affect Scores. However, I also subtracted the mean and divided by the standard deviation of the Democratic Affect Scores and Republican Affect Scores during model estimation, as recommended by [Clinton, Jackman and Rivers \(2004\)](#), to ensure identifiability. Having recovered estimates of affect toward each of the parties for each politician’s Twitter account, I now proceed to validate the output of my framework for measuring affective polarization of politicians.

Description and Face Validity

First, I show that the parties in government do appear to be affectively polarized using my measures. In [Figure 4](#), we can see that Democrats have largely positive Democratic Affect Scores and Republicans have mostly negative Democratic Affect Scores. This indicates that Democrats are more positive toward Democrats and that Republicans are more negative toward Democrats. We can see two distinct modes and relatively little overlap between

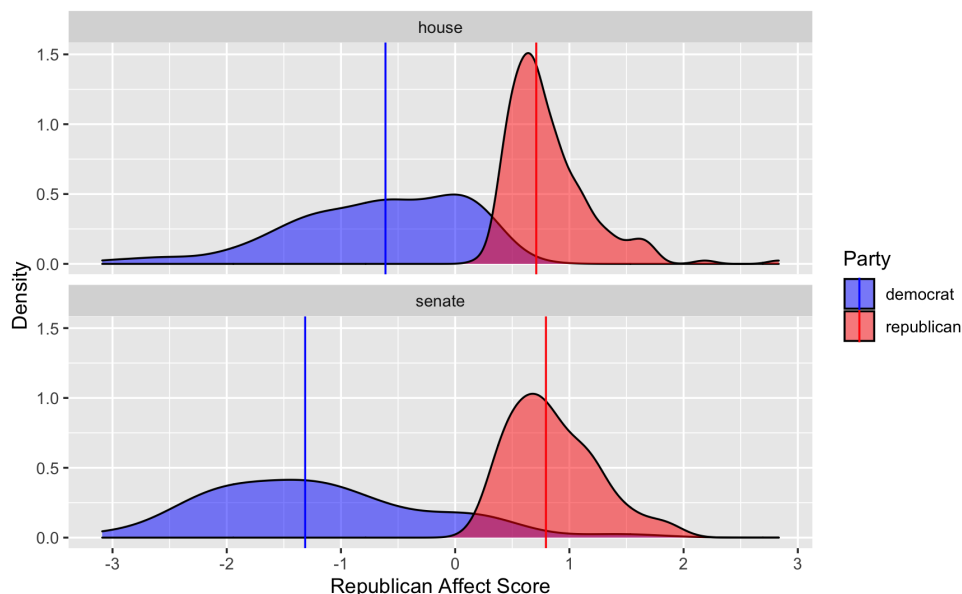
Figure 4. Democratic Affect Score Distributions in the House and Senate



them—as we would expect given affectively polarized elites. Moreover, the party medians, indicated by the vertical lines, are quite far from one another. In [Figure 5](#), I display the analogous density plots for Republican Affect Scores in the House and Senate. Here, we see that Democrats are much more negative toward Republicans in their tweets than are Republicans. Interestingly, the distributions of out-party affect scores seem to be somewhat flatter, indicating a wide range of views about the opposing party, while the in-party affect distributions are more narrow, indicating less differentiation with respect to in-party affect.

Turning to a more fine-grained analysis, we can see that the arrangements within each party largely make sense. Since this measure captures both messaging and affective polarization, there are two important groups this measure should identify near the poles: party leaders and individuals who are widely reputed to be intensely partisan. As [Lee \(2016\)](#) notes, party leaders should devote a large amount of time to cultivating their party’s brand through both positive in-party messaging and negative messages about the opposing party. Further, some rank-and-file MCs have developed reputations as strident defenders of their own party and in terms of overwhelmingly negative attacks against the opposing party (see, for example, [Breshnahan, Zanona and Cheney 2020](#) and [Zanona 2019](#)).

Figure 5. Republican Affect Score Distributions in the House and Senate



In [Figure 6](#), I plot Republican Affect Scores against Democratic Affect Scores for the members of the House. In the upper left quadrant of the plot are members who are extremely negative toward Democrats and extremely positive toward Republicans. The lawmakers identified in this section are widely reputed to be intensely partisan. For example, Reps. Doug Collins (R-GA), Matt Gaetz (R-FL), Andy Biggs (R-AZ), Jim Jordan (R-OH), and Paul Gosar (R-AZ) are all among the most steadfast Trump allies in Congress, attacking Democrats and defending the Republican Party. The House Republican leadership is also in the upper left of the plot: Minority Leader Kevin McCarthy (R-CA) and Minority Whip Steve Scalise (R-LA), who are incentivized to tarnish the brand of their opponents and to bolster the brand of their own party ([Lee 2016](#)), have extremely negative Democratic Affect Scores and extremely positive Republican Affect Scores. The arrangement for Democrats also makes intuitive sense: Speaker Nancy Pelosi (D-CA) and Majority Leader Steny Hoyer (D-MD) are both clustered in the lower right of the plot, meaning they have high Democratic Affect Scores and low Republican Affect Scores. Representatives from both parties who have reputations for working with their partisan opponents—and who, presumably are less affectively polarized—are clustered around the center of the plot (e.g., Reps. Anthony

Brindisi (D-NY) and Angie Craig (D-MN)).

Figure 6. Partisan Affect Scores in the 116th House



Figure 7. Partisan Affect Scores in the 116th Senate



Figure 7 displays the same plot for members of the 116th U.S. Senate. Many of the most extreme partisans in the Republican Party are clustered in the upper left quadrant, indicating that they have the highest Republican Affect Scores and the lowest Democratic Affect Scores. Sens. Ted Cruz (R-TX), Marsha Blackburn (R-TN), and Rick Scott (R-FL) have all cultivated images of intense partisanship, including attacking partisan opponents and stridently defending their own party. Moving toward the origin, where senators who are not extreme in either Democratic or Republican Affect Scores are located, we can see notable bipartisan senators such as Sens. Lisa Murkowski (R-AK), Joe Manchin (D-WV), and Kirsten Sinema (D-AZ). These are the senators we would expect to be less affectively polarized as they regularly work across party lines, and Manchin and Sinema have regularly thwarted Democratic plans during the 117th Congress (Bade et al. 2021). Finally, in the lower right quadrant, Democratic leadership appear to be the most positive toward the Democratic Party and the most negative toward the Republican Party. Sens. Chuck Schumer (D-NY) and Richard Durbin (D-IL) have extremely high Democratic Affect Scores and extremely low Republican Affect Scores. There are some questionable locations on the Republican side. For example, Majority Leader Mitch McConnell (R-KY) appears to be a middle-of-the-road Republican senator in terms of his Democratic Affect Score and is quite low on his Republican Affect Score—despite his position as a party leader. This could be due to the use of his official Twitter account, which, in this case, was a “press office” account as was the account for Sen. Josh Hawley (R-MO). These accounts appear to be run by staff and often post tweets more like press releases while unofficial, personal accounts appear to be more partisan. McConnell also was less Trump-aligned than many of his colleagues and voted recently for the bipartisan infrastructure bill under President Biden.

In sum, it appears that both Democratic and Republican Affect Scores divide the parties clearly and intuitively, identifying party leadership and intensely partisan members of the United States House of Representatives and Senate. Similarly, it does a fairly good job singling out the more bipartisan, moderate members of Congress, such as those willing

to cross party lines on key votes like impeachment.

As a more systematic analysis, [Table 4](#) displays comparisons of median affect scores for several House caucuses as well as the overall party medians. The House Freedom Caucus has earned a reputation for intense partisan combativeness, with Reps. Jordan, Gosar, and Marjorie Taylor Greene (R-GA) as members. Particularly during the first impeachment of then-President Trump, members of the House Freedom Caucus were some of his most aggressive defenders, regularly attacking Democrats. Affect scores reveal this tendency as members of the House Freedom Caucus are both more negative toward Democrats and more positive toward Republicans than the Republican medians—precisely what we would expect from a measure of partisan affect. Though there is not a neat analogue of the House Freedom Caucus on the Democratic side, the House Progressive Caucus has taken a harder line against the Republican Party, preferring to forgo bipartisan compromise. Indeed, the Progressive Caucus’s median affect scores are more extreme than the Democratic Party medians. Finally, there are several groups of lawmakers in the House who bill themselves as bipartisan, and in fact, the House Problem Solvers Caucus is bipartisan. In all cases, these bipartisan caucuses have less extreme affect scores than their party medians. Thus, MCs who are more willing to work with members of the opposing party, and therefore, logically should be less affectively polarized, are rated as such by my measures. Moreover, MCs demonstrate their commitment to intense partisanship through their caucus memberships tend as well to score as more extreme than their parties on my measures.

These preliminary findings demonstrate that my framework for measuring partisan affective polarization using tweets is broadly consistent with expectations. Partisan “bomb-throwers,” identified as such by [Zanona \(2019\)](#), and party leaders are located near the extremes of Democratic and Republican Affect Score distributions, and moderate, bipartisan members of both parties are similarly correctly situated in the overlapping, bipartisan center of the distributions. While these measures appear to have broad face validity, further

Table 4. Comparison of Party and Caucus Affect Score Medians

	Democratic Party				
	Party Median	Progressive Med.	Problem Solvers Med.	Blue Dog Med.	New Dem. Med.
Democratic Affect Score	0.67	1.03	0.27	0.25	0.56
Republican Affect Score	-0.61	-0.93	0.04	0.14	-0.40
	Republican Party				
	Party Median	Freedom Caucus Med.	Problem Solvers Med.	Tuesday Group Med.	
Democratic Affect Score	-0.67	-1.12	-0.25	-0.35	
Republican Affect Score	0.71	0.85	0.62	0.67	

analyses are needed to confirm that I am capturing affective polarization and messaging and not something else. In the next section, I attempt to carry out just such analyses to further validate this measure.

Validation

If affect scores are capturing the underlying affective polarization and prioritization of messaging of MCs, we can draw several empirical predictions from the literature about how they should relate to behavior. In this section, I test whether affect scores are related to behaviors as predicted. Specifically, I examine whether MCs who are affectively polarized in messaging according to my measures are less bipartisan and more strident partisan defenders on other measures than the less affectively polarized. Further, I examine the relationship between mass-level affective polarization and job approval of more extreme partisan messagers in Congress.

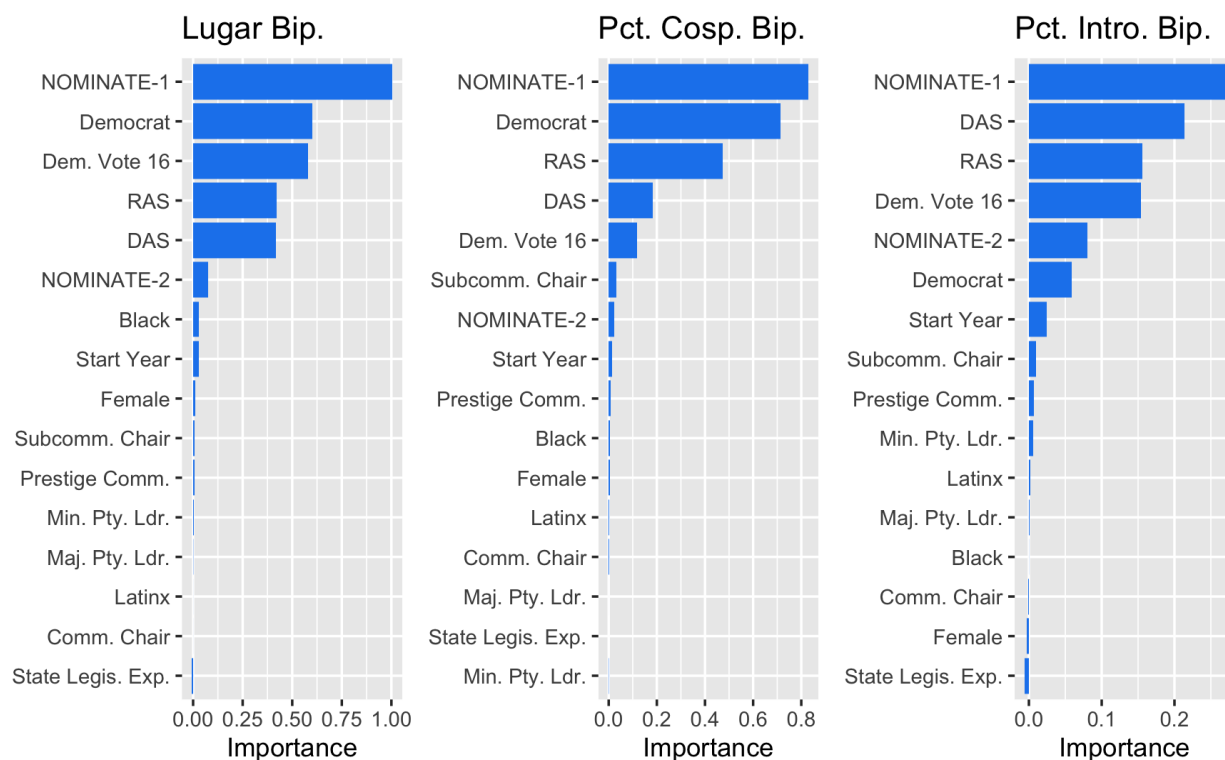
Bipartisanship

Logically, sending intensely negative messages about the opposing party is not conducive to fostering close working relationships across party lines ([Lee 2016](#)). Further, affective polarization, by definition, implies a relative distaste for members of the other party. I assume, therefore, that MCs who prioritize affectively charged party messaging should be less bipartisan in their behavior. I draw measures of bipartisanship from several sources: the Lugar Center calculates weighted indexes of bipartisan behavior ([The Lugar Center and the McCourt School of Public Policy 2021](#)). Govtrack measures bipartisanship using the percent of bills cosponsored by an MC that are originally sponsored by a member of the other party and the number of bills introduced by an MC with bipartisan cosponsors ([Govtrack 2021](#)). To aid in presentation, I divide the number of bills with bipartisan cosponsors by the number of introduced bills and multiply by 100, leaving the percent of introduced bills with bipartisan cosponsors. With all of these measures, bill cosponsorship is taken as an indicator of bipartisanship.

Random forests provide a good test of the relationship between bipartisanship and my tweet-based measures of affect. Due to the randomness in variable selection and observations, random forests allow for modeling complex interactions between variables, essentially taking functional form specification out of the hands of the researcher. Moreover, during fitting of the random forest models, the algorithm randomly permutes the values of each variable and predicts outcomes for out-of-bag observations. The increase in mean-squared error of predictions when a variable's values are randomly permuted provides an estimate of the importance of that variable to prediction ([Breiman 2001](#)). For these reasons, I fit random forests with 10,000 trees, predicting each of the measures of bipartisanship in the House and Senate. In addition to Democratic and Republican Affect Scores, I include gender, race, seniority, leadership positions, Democratic presidential vote share in 2016, membership on "prestige" or "power" committees, and, importantly, party as well as both dimensions of DW-NOMINATE. [Harbridge \(2015\)](#) and [Volden and Wiseman \(2018\)](#) cite each of these

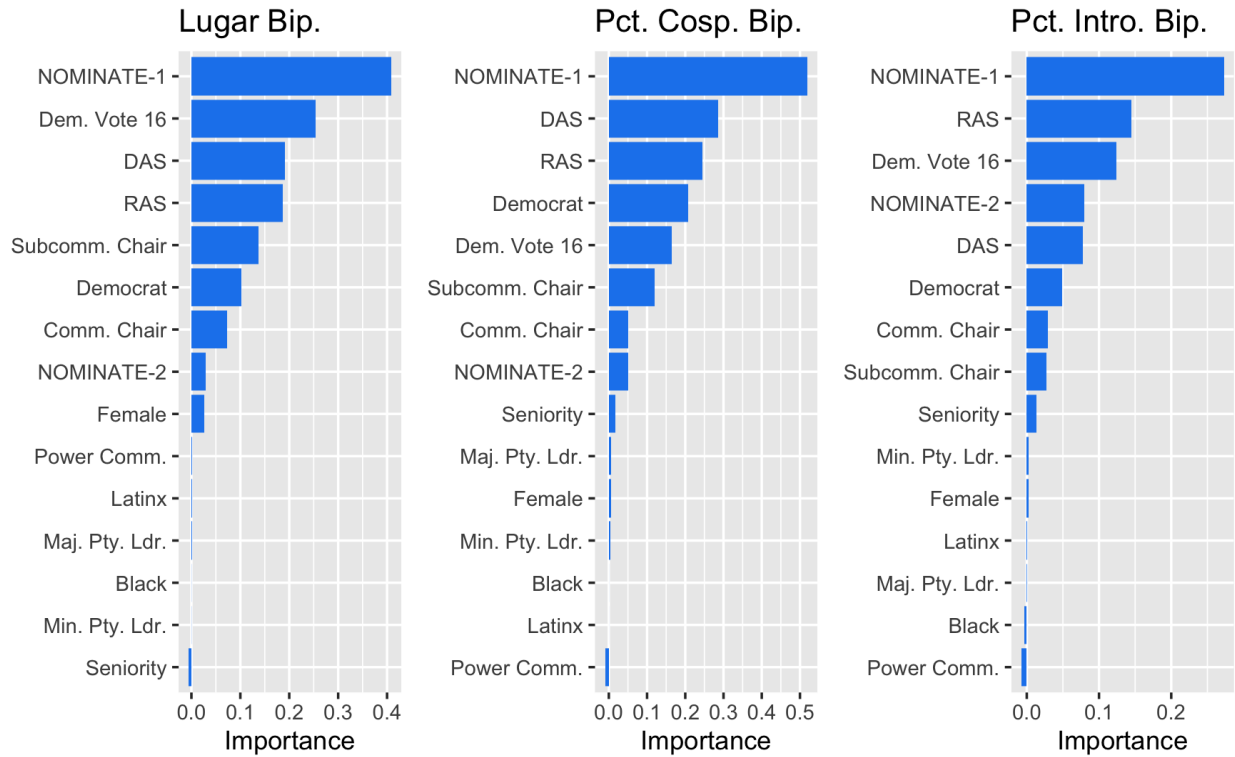
variables as key predictors of bipartisanship and legislative productivity in the House and Senate. Figure 8 displays variable importance plots for random forests fit in the House of Representatives, and Figure 9 presents these permutation importance plots from the Senate.

Figure 8. Random Forest Variable Importance Plots for House Bipartisanship



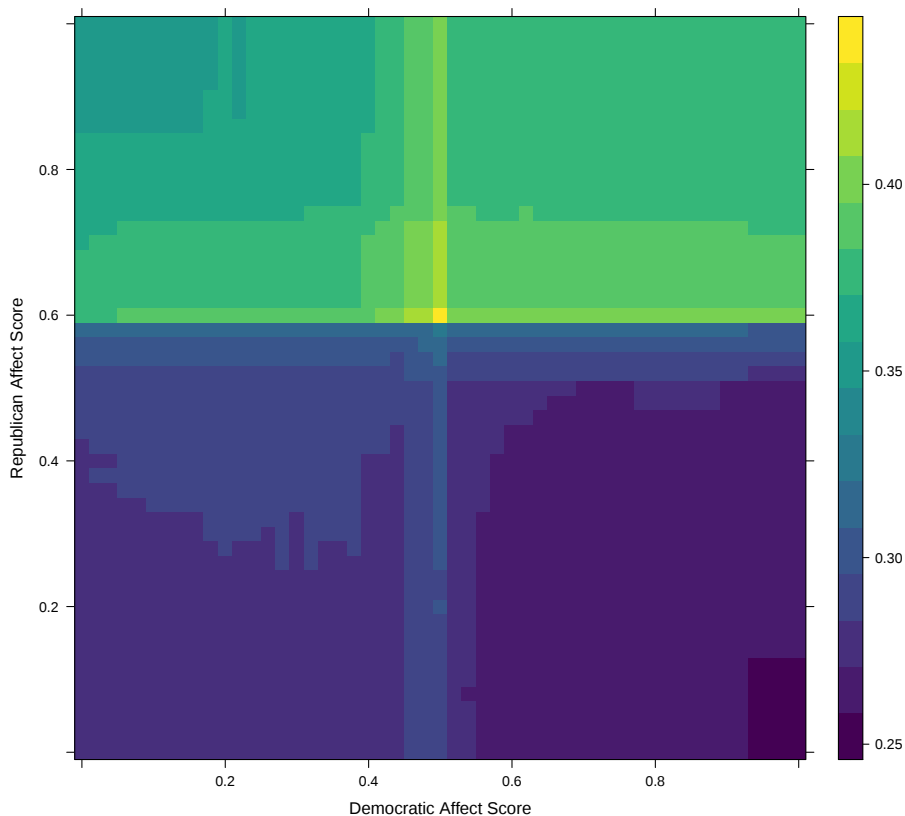
While the first dimension of DW-NOMINATE ranks as the most important predictor of every measure of bipartisanship, Democratic and Republican Affect Scores regularly rank among the most important predictors of bipartisanship—in many cases, rivalling or outstripping constituency partisanship and the legislators’ own partisanship in importance. That affect scores come close to the importance of first dimension NOMINATE is also striking since NOMINATE largely measures the partisanship of a member’s voting record (Poole and Rosenthal 1997). Briefly, then, it appears that affect scores are highly predictive of member bipartisanship, just as we would expect if affect scores were capturing the affective polarization of members. Members who are more affectively polarized should be less willing

Figure 9. Random Forest Variable Importance Plots for Senate Bipartisanship



to work with their partisan opponents on legislation, even after accounting for ideology and constituency influence. This relationship is shown in [Figure 10](#), which displays a partial dependence plot of predicted standardized cosponsorship proportions against affect scores from a random forest with 500 trees. In [Figure 10](#), it is clear that the most affectively polarized Democrats (those in the bottom right quadrant) and Republicans (upper left quadrant) have lower rates of bipartisan cosponsorships while the least affectively polarized members of the House have higher rates of bipartisan behavior. Moving from the highest Democratic Affect Score and lowest Republican Affect Score to the least polarized middle of the affect score distribution, the proportion of cosponsorships which are bipartisan is predicted to increase from about 0.25 to 0.42 (an increase of 17 percentage points). Similarly, moving from the highest Republican Affect Score and lowest Democratic Affect Score to the least polarized middle of the distribution, the proportion of cosponsorships which are bipartisan is expected

Figure 10. Predicted Bipartisan Proportion of Cosponsorships (House)



to increase from about 0.35 to 0.42 (an increase of 7 percentage points). These predicted shifts in bipartisan behavior are substantively large, representing an increase of nearly one standard deviation and a third of a standard deviation, respectively. Clearly, then, affect scores are important predictors of bipartisan behavior—even after taking party, NOMINATE scores, and electoral considerations into account.

Trump Loyalty

Particularly during the Trump era, partisan rancor in government seems to have reached a new peak (Fox 2021; Itkowitz and DeBonis 2021; Kaplan 2021). Many Republican MCs have built their reputations around strict support for Trump and intensely negative messaging against Democrats—in effect, styling themselves after Trump himself. Again, if my framework for measuring partisan affective polarization and messaging is truly capturing the

intended constructs, we would expect these intensely partisan loyalists to stand out in their affect scores. Those who identify most strongly with their party should be most likely to defend their party—even in the face of controversy. Likewise, those who hate the opposing party should be the most likely to attack the opposing party.

Axios created a measure of Trump loyalty, assigning members a score for their response to each of seven controversies during Trump’s time in politics (Bartz et al. 2020). Members who most vocally supported Trump during each of these scandals were given higher scores. Additionally, Axios averaged these loyalty scores with the percent of votes taken in alignment with Trump’s preferences. Republican members of Congress were given Trump loyalty scores that ranged from a minimum of 0 to a maximum of 100—though the observed minimum and maximum were 27 and 99, respectively. I utilize these scores to examine whether IMAP matches the subjective perceptions of political observers in the news media and Trump-era voting records. In this case, a vocal defense of Trump, as categorized by Axios, reflects, at a minimum, a positive affect toward the Republican Party and likely, a negative affect toward the Democratic Party as well. Importantly, Trump loyalists in Congress also tended to be those who most harshly attacked Democrats. In all, then, Trump loyalty should serve as another measure of the affective polarization of Republican MCs against which to benchmark politicians’ Affect Scores.

Table 5 shows results from a regression of Axios Trump loyalty scores on Democratic and Republican Affect Scores, adjusting for both dimensions of DW-NOMINATE (Lewis et al. 2021) and Democratic vote share in the 2016 election. I also include models with Bonica’s (2014) CFScores and Barber and McCarty’s (2015) “Tweet Scores.” I include NOMINATE, CFScores, and TweetScores to account for the possibility that Trump loyalty scores are influenced by conservatism or by MC’s voting records. Variables are standardized, and heteroskedasticity-robust standard errors are shown in parentheses. Even after controlling for ideology and constituency influence, out-party (Democratic) Affect Score is a significant predictor of loyalty to Trump. In the House, a standard deviation increase in Democratic

Affect Score, indicating relatively more positive affect toward Democrats, is associated with between one fifth and one fourth of a standard deviation decrease in Trump loyalty. A standard deviation increase in Democratic affect in the Senate (i.e., a standard deviation increase in relative positivity toward Democrats) is associated with a nearly one half standard deviation decrease in Trump loyalty. Inversely, Republicans who are more negative toward the Democratic Party in their affect scores are expected to be more loyal to Trump (i.e., more stridently defensive of Trump through major scandals). In five of the six regressions presented in [Table 5](#), adjusting for ideology using NOMINATE, TweetScores, or CFScores, Democratic Affect Score is a significant predictor of Trump loyalty. Interestingly, it seems that in-party (Republican) Affect Scores do not approach statistical significance. That said, the direction of these coefficients indicates that politicians who are more positive toward the Republican Party in their tweets are more loyal to Trump. The results in this section indicate that, as expected, Republicans who are rated as more loyal to Trump tend to be more negative toward the Democratic Party, according to their Democratic Affect Scores. To put it somewhat differently, the perceptions of political observers in the media align with my measure quite well, indicating that my framework successfully captures the intended construct of intense partisanship among Republicans. In the next section, I show that affect scores also capture some of the variation in activist perceptions of politicians.

Activist Perceptions of Conservatism

As discussed in the previous section, loyalty to Trump, primarily indicated by vociferous defenses of the Republican Party and denunciations of the Democrats, has become synonymous with conservatism, especially among partisan activists. Republican MCs have been censured and challenged in primaries by local party activists for simply voting against Trump during his second impeachment ([Ruwitch and Sprunt 2021](#); [Slotkin 2021](#)). Additionally, Rep. Liz Cheney (R-WY), despite being quite conservative ideologically, was replaced in her leadership position for contradicting Trump’s claims of election fraud by Rep. Elise Stefanik (R-NY), who is a stronger defender of Trump ([Sprunt 2021](#)). These recent developments

Table 5. Affect Scores and Trump Loyalty Among Republicans

	Trump Loyalty Score					
	House			Senate		
	(1)	(2)	(3)	(4)	(5)	(6)
Republican Affect Score	0.086 (0.119)	0.155 (0.115)	0.158 (0.117)	0.196 (0.189)	0.144 (0.166)	0.125 (0.158)
Democratic Affect Score	-0.229* (0.119)	-0.254** (0.110)	-0.284** (0.112)	-0.385** (0.192)	-0.421** (0.183)	-0.468*** (0.168)
NOMINATE - Dim. 1	0.356*** (0.089)			0.195 (0.141)		
NOMINATE - Dim. 2	0.165** (0.083)			0.217 (0.169)		
CFScore		0.902** (0.374)			0.670 (0.813)	
Tweetscore			-0.105 (0.195)			-0.015 (0.399)
Democratic Vote 2016	-0.016 (0.011)	-0.024** (0.012)	-0.027** (0.012)	-0.021 (0.017)	-0.039** (0.017)	-0.039** (0.017)
Constant	0.576 (0.364)	-0.056 (0.520)	1.070** (0.444)	0.821 (0.672)	0.840 (0.971)	1.555** (0.683)
Observations	172	171	173	48	45	47
Adjusted R ²	0.336	0.299	0.261	0.376	0.278	0.315

Note:

*p<0.1; **p<0.05; ***p<0.01

indicate that, to Republican activists in the Trump era, “to be conservative is partly to support Trump” (Hopkins and Noel 2022, p. 1). Hopkins and Noel (2022) show systematic evidence that activists indeed perceive Trump loyalists in the 114th Senate to be significantly more conservative than Republicans who opposed Trump. Extending these results, if activist perceptions of conservatism are indeed influenced by Trump loyalty, affect scores should be significantly related to perceived ideology. Below, I test this relationship for Senate Republicans who served in the 114th Congress, using estimates of ideology derived from activists’

pairwise comparisons of senators by [Hopkins and Noel \(2022\)](#).

[Table 6](#) displays the results of several regressions of pairwise estimates of ideology on DW-NOMINATE as well as [Hopkins and Noel’s \(2022\)](#) measure of Trump support. Again, standard errors are heteroskedasticity-robust, and variables are standardized. As we can see in the table, Republican Affect Scores are a highly significant predictor of activist perceptions of conservatism: Republicans who are more positive in their affect toward the Republican Party are far more likely to be perceived as conservative. An increase of one standard deviation in Republican Affect Score is associated with an increase of more than one half of a standard deviation in perceived conservatism. These coefficient estimates come quite close to the coefficient estimates for NOMINATE’s first dimension, indicating again, that affect scores are quite important predictors. Democratic Affect Scores fail to obtain statistical significance at traditional levels in these regressions, though their coefficient estimates indicate that Republican senators who are more positive toward Democrats are perceived as less conservative than are more negative Republican senators.

Overall, then, the results of the last two sections show that affect scores are significantly related to the perceptions of various political observers, including members of the media and partisan activists. Specifically, media perceptions of Trump loyalty among Republicans are significantly predicted by affect scores. Trump loyalty, characterized primarily by strident defenses of the Republican Party and attacks against Democrats during notable Trump scandals, should match up quite well with a measure of partisan animosity, and in fact, affect scores do match up quite well with Trump loyalty. Further, partisan activists’ perceptions of conservatism, which has become more or less synonymous with intense partisanship since Trump’s first candidacy in the 2016 election, align with affect scores as well. In both cases, Republicans identified as the most vocal partisan defenders and out-partisan antagonists are identified as more extreme by my framework. These overlaps between media and activist perceptions indicate that affect scores are valuable measures of partisan animus among elites.

Table 6. Affect Scores and Activist Perceptions of Conservatism among Senate Republicans

	<i>Dependent variable:</i>		
	Pairwise Activist Ideology		
	(1)	(2)	(3)
NOMINATE-1	0.826*** (0.114)	0.592*** (0.101)	0.605*** (0.103)
NOMINATE-2	0.141* (0.073)	0.080 (0.059)	0.091 (0.061)
Anti-Trump	-0.052 (0.136)		0.092 (0.115)
DAS		-0.152 (0.094)	-0.172* (0.098)
RAS		0.564*** (0.155)	0.576*** (0.156)
Constant	-0.010 (0.126)	-0.312*** (0.110)	-0.361*** (0.126)
Observations	41	41	41
Adjusted R ²	0.614	0.753	0.751
<i>Note:</i> *p<0.1; **p<0.05; ***p<0.01			

Approval of Representatives and Affective Polarization

Though there is no direct measure of elite-level affective polarization to which I can compare affect scores, there is a consistent measure of affective polarization among the mass public which presumably would be related to elite-level affective polarization. Feeling thermometers capture how warmly (or coldly) survey respondents feel about the two major parties on a scale from 0 to 100. It seems reasonable to assume that more affectively polarized members of the public should be more likely to approve of their representatives in Congress when those representatives are themselves affectively polarized—especially, if that affective polarization manifests itself through the denigration of the opposing party and the defense of the in-party. People who hold more extreme feelings toward the parties may approve more of

politicians who express these views as well because they feel represented by these politicians. Additionally, members of the public might feel that politicians who consistently message negatively about one party and positively about the other are defending (or attacking) their partisan identity or even, that the politicians have similar identities.

The American National Election Study’s 2020 survey includes both feeling thermometers and measures of approval of the House incumbent in the respondent’s district ([ANES 2021](#)). This allows me to study the relationship between affect scores, mass-level affective polarization, and approval for House incumbents. Whether due to representational or identity-based considerations, if affect scores successfully capture partisan messaging and elite affect toward the parties, this should be reflected in public approval—especially among members of the public with extreme affective reactions toward the two parties. In other words, there should be an interaction between affect scores and feeling thermometer ratings of the parties in predicting approval of House incumbents.

In [Table 7](#), results are displayed from several multilevel models with state random effects, predicting binary approval of House incumbents. Each model adds additional respondent- and representative-level control variables, including party identification and ideology and the interactions between incumbent and respondent variables. Importantly for my results, the interaction terms between representative Republican Affect Scores and feeling thermometer ratings of the parties are strong and significant, indicating that members of the public who feel more negatively toward Democrats are more likely to approve of representatives who express positive affect toward Republicans in their tweets. Similarly, members of the public are more likely to approve of representatives with high Republican Affect Scores if they themselves feel more positively toward the Republican Party. In these results, there is a clear relationship between elite affect in tweets and mass affect toward the parties from feeling thermometer ratings. That said, the coefficients for interactions between Democratic Affect Scores and party feeling thermometer ratings are insignificant—though they are in the expected directions: respondents who like the Democratic Party are more likely to approve

of representatives with more positive Democratic affect, and those who like the Republican Party more are less likely to approve of representatives who express positive affect about the Democratic Party).

Table 7. Affective Polarization and Approval of House Incumbents

	<i>Dependent variable:</i>		
	Approval of House Incumbent		
	(1)	(2)	(3)
Dem. Feeling Therm.	0.348*** (0.054)	0.351*** (0.054)	0.380*** (0.064)
Rep. Feeling Therm.	0.169*** (0.052)	0.170*** (0.052)	0.101* (0.061)
DAS	0.100 (0.072)	0.044 (0.074)	0.048 (0.083)
RAS	-0.040 (0.075)	-0.003 (0.076)	0.059 (0.086)
DAS*Dem. Feeling Therm.	0.101 (0.087)	0.093 (0.087)	0.110 (0.099)
DAS*Rep. Feeling Therm.	-0.159* (0.088)	-0.154* (0.088)	-0.091 (0.102)
RAS*Dem. Feeling Therm.	-0.335*** (0.090)	-0.339*** (0.090)	-0.242** (0.103)
RAS*Rep. Feeling Therm.	0.248*** (0.087)	0.249*** (0.088)	0.299*** (0.104)
Constant	✓	✓	✓
Respondent Party ID	✓	✓	✓
Incumbent Party	✓	✓	✓
Inc. Party * Respondent Party	✓	✓	✓
NOMINATE-1		✓	✓
Respondent Ideo. ID			✓
NOMINATE * Resp. Ideo.			✓
Observations	6,134	6,122	5,287
<i>Note:</i> *p<0.1; **p<0.05; ***p<0.01			

In short, it appears that affective polarization at the mass level predicts favorabil-

ity toward politicians who are more extreme in their messaging about the parties: people who feel more positively toward the Republican Party are more likely to feel warmly about politicians who express more positive views about Republicans. The next section considers whether these results extend beyond the 116th Congress to U.S. governors.

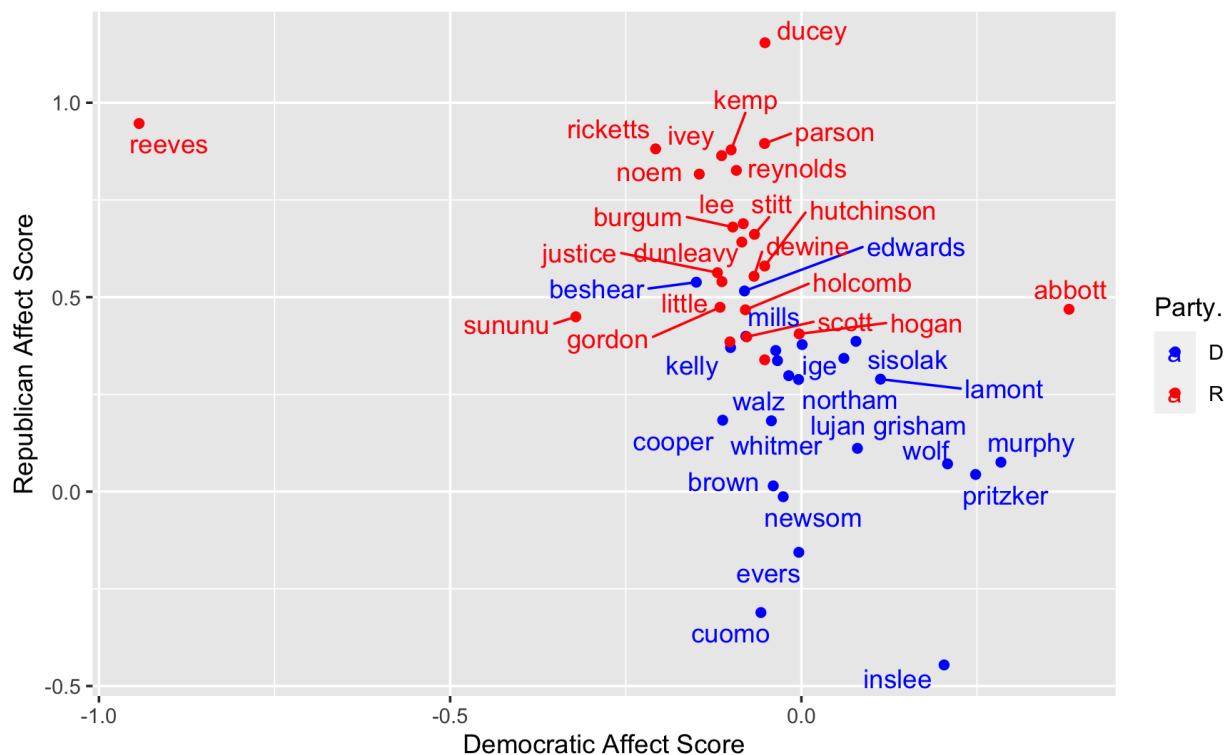
State Governors

As a final check of the validity of affect scores, I test whether the above findings in Congress hold in a different sample of politicians. Namely, whether the affect scores of American governors display similar face validity and whether they relate to mass-level partisan affect. [Figure 11](#) plots gubernatorial Republican Affect Scores against Democratic Affect Scores.

The ordering of governors largely makes intuitive sense given the reputations of the individual governors. Most of the extremely partisan Republican governors have high Republican Affect Scores. For instance, Govs. Kemp (GA), Noem (SD), and Ducey (AZ) all cluster near the top in terms of their Republican Affect Scores. Meanwhile, more bipartisan governors, like Govs. Dewine (OH) and Hogan (MD) appear near the middle of the plot. On the Democratic side, governors who have reputations for moderation like Govs. Beshear (KY) and Edwards (LA) are located closer to the Republicans in their affect scores. Notably, Gov. Abbott of Texas appears to have a strikingly high Democratic Affect Score, a counter-intuitive result. Indeed, a manual examination of Abbott’s Twitter indicates that several of his tweets were misclassified as mentioning the Democratic Party due to the acronym for the Division of Emergency Management (“DEM”). This can be fixed by including more tweets from governors in the coding process and perhaps by fitting separate prediction models for governors. Still, overall, the arrangement of governors appears to make sense. Moreover, the affect scores among governors are less extreme than those in Congress, consistent with the findings of [Kousser and Phillips \(2012\)](#) that governors are less partisan and more willing to work with partisan opponents.

Finally, I replicate my results for the approval of House incumbents using the sample of governors and the 2020 ANES. Unfortunately, the 2020 ANES did not ask a standard

Figure 11. Partisan Affect Scores Among Governors



job approval question for governors but asked specifically about approval of each governor's COVID response (ANES 2021). Consequently, results in this section should be interpreted somewhat tentatively: the relationship between approval and party feeling thermometer ratings could be due to the feeling of identity representation from governors' affective styles, or it could be due to real differences in COVID policies between parties. Despite these caveats, I find that members of the public who are more affectively polarized are more likely to approve of governors who express similar affect in their tweets. In Table 8, I fit a multilevel logistic regression with random effects for states. Adjusting for party identification of the respondent and partisanship of the governor, these results indicate that respondents who are more positive toward the Republican Party on feeling thermometer measures are more likely to approve of the COVID response of governors who are more positive toward the Republican Party (and negative toward the Democratic Party) in their tweets.

Contrary to expectations, respondents who are more positive toward Democrats are less likely to approve of governors who express more positive affect toward Democrats, and respondents who are more positive toward Republicans are more likely to approve of governors who express more positive affect toward Democrats. After some investigation, it appears that this was largely due to Texas Governor Abbott whose tweets were misclassified, as discussed above. Model 2 of [Table 8](#) omits respondents from Texas, confirming that the counterintuitive significant results were due to the misclassification of Abbott’s tweets. In the future, I hope to generate separate models for governors and state legislatures by coding larger samples of tweets from governors and state representatives specifically which could avoid such classification errors. However, generally, it appears that the findings from Congress apply fairly well to state governors. The arrangements of governors mostly make sense on their face, and mass-level affect appears significantly related to the affect expressed in governors’ tweets.

Conclusion and Discussion

In recent years, there has been a rise in the temperature of elite discourse—highlighted especially by the Trump presidency. However, the discipline of political science has not yet developed a firm understanding of this discourse or the role it plays in our political system due largely to a lack of systematic measurement. In this paper, I developed and validated a framework for measuring affective polarization in elite messaging, IMAP, intended to capture the increasingly partisan discourse of political elites. I employed natural language processing to categorize a large number of tweets and fit a Bayesian latent variable model to recover the underlying features of the Twitter data I collected from politicians. First, I demonstrated the face validity of partisan Affect Scores from this framework, showing that intensely partisan political figures and party leaders were identified as such by my measure. Additionally, this measure also correctly placed bipartisan politicians. Then, I showed that affect scores predicted bipartisan cosponsorship behavior and Trump loyalty—even after

Table 8. Affective Polarization and Approval of Gubernatorial COVID Response

	<i>Dependent variable:</i>	
	Approval of Governor's COVID Response	
	(1)	(2)
Dem. Feeling Therm.	0.623*** (0.048)	0.710*** (0.050)
Rep. Feeling Therm.	0.034 (0.046)	-0.071 (0.049)
DAS	-0.074 (0.097)	0.031 (0.092)
RAS	-0.393*** (0.142)	-0.389*** (0.144)
DAS*Dem. Feeling Therm.	-0.120*** (0.035)	-0.028 (0.043)
DAS* Rep. Feeling Therm.	0.115*** (0.037)	-0.079 (0.048)
RAS*Dem. Feeling Therm.	-0.448*** (0.046)	-0.424*** (0.050)
RAS*Rep. Feeling Therm.	0.375*** (0.045)	0.306*** (0.048)
Constant	✓	✓
Respondent Party ID	✓	✓
Incumbent Party	✓	✓
Inc. Party*Respondent Party	✓	✓
Omit Governor Abbott (Texas)		✓
Observations	7,415	6,832

Note: *p<0.1; **p<0.05; ***p<0.01

accounting for traditional measures of ideology and constituency influence. Finally, I showed that American governors are also polarized on these measures, though less so than members of Congress. Importantly, affect scores also appear to be related to mass-level affective polarization: interactions between party feeling thermometer ratings from ANES and IMAP were significant predictors of approval for both House incumbents and governors.

Given recent political violence that was encouraged in large part by prominent political figures, it is important that we are able to systematically measure how affectively polarized elites are in their communications. Not only can this help us understand elite behavior, but it can, hopefully, help us understand the dynamics of elite rhetoric and mass behavior. [Mutz and Reeves \(2005\)](#) finds that exposure to uncivil elite communications can reduce trust in government. My preliminary results suggest that elites are quite uncivil in their communications about members of the other party. Future work could employ IMAP to understand how this relates to public opinion, to public trust in government, and potentially, to political violence.

With a more systematic measure in hand, future work can examine how affective polarization changes dynamically, how the public has affectively polarized, and how elite and mass affective polarization interact with one another. Surveys of elites could also be used to examine whether this measure relates to traditional measures of affective polarization like party feeling thermometers. This measure could be applied to state legislators and local governments to examine variation across levels of government. These measures could also be used to place politicians and members of the public on the same scale. In short, the study of affective polarization and messaging is a potentially fruitful avenue for future research, to which I believe my framework presents an important first step.

Bibliography

- ANES. 2021. “ANES 2020 Time Series Study Full Release [dataset and documentation]. July 19, 2021 version.”
URL: www.electionstudies.org
- Ansolabehere, Stephen, James Snyder and Charles Stewart. 2001. “Candidate Positioning in U.S. House Elections.”
- Bade, Rachael, Eugene Daniels, Eli Okun and Garrett Ross. 2021. “Politico Playbook PM: The Problem with Manchin/Sinema and Voting Rights.” *Politico* .
URL: <https://www.politico.com/newsletters/playbook-pm/2021/06/03/the-problem-with-manchin-sinema-and-voting-rights-493108>
- Bafumi, Joseph and Michael Herron. 2010. “Leapfrog Representation and Extremism: A

- Study of American Voters and Their Members in Congress.” *American Political Science Review* 104(3):519–542.
- Barber, Michael and Nolan McCarty. 2015. Causes and Consequences of Polarization. In *Solutions to Political Polarization in America*, ed. Nathan Persily. New York: Cambridge University Press.
- Bartz, Juliet, Mike Allen, Jim VandeHei and Orion Rummeler. 2020. “Always Trumpers: The President’s Unbreakable Wall.” *Axios* .
URL: <https://www.axios.com/always-trumpers-republicans-axios-on-hbo-f4a866a6-03a8-4632-b92b-47f79607ef01.html>
- Bonica, Adam. 2014. “Mapping the Ideological Marketplace.” *American Journal of Political Science* 58(2).
- Breiman, Leo. 2001. “Random Forests.” *Machine Learning* 35:5–32.
- Breshnahan, John, Melanie Zanona and Kyle Cheney. 2020. “McCarthy Embraces Ex-Rival Jordan as the Top Partisan Fighter.”
URL: <https://www.politico.com/news/2020/05/08/kevin-mccarthy-embraces-jim-jordan-243508>
- Cineas, Fabiola. 2021. “Donald Trump is the Accelerant.” *Vox* .
URL: <https://www.vox.com/21506029/trump-violence-tweets-racist-hate-speech>
- Clinton, Joshua, Simon Jackman and Douglas Rivers. 2004. “The Statistical Analysis of Roll-Call Data.” *American Political Science Review* 98(2):355–370.
- Cox, Gary and Mathew McCubbins. 1993. *Legislative Leviathan: Party Government in the House*. Berkeley: University of California Press.
- Cox, Gary and Mathew McCubbins. 2005. *Setting the Agenda: Responsible Party Government in the U.S. House of Representatives*. Cambridge University Press.
- Edmondson, Catie. 2021. “House Republicans Who Backed Infrastructure Bill Face Vicious Backlash.” *The New York Times* .
URL: <https://www.nytimes.com/2021/11/10/us/politics/republicans-backlash-infrastructure-bill.html>
- Fenno, Richard. 1977. “U.S. House Members in Their Constituencies: An Exploration.” *The American Political Science Review* 71(3):883–917.
- Fenno, Richard. 1978. *Home Style: House Members in Their Districts*. Boston: Little Brown and Company.
- Fox, Lauren. 2021. “‘Toxic Is Spot-On’: House Members Describe Roiling Animosity Among Lawmakers.”
URL: <https://www.cnn.com/2021/11/23/politics/congress-anger-house-fights/index.html>

- Gelman, Andrew and Jennifer Hill. 2007. *Data Analysis Using Regression and Multi-level/Hierarchical Models*. Cambridge University Press.
- Gift, Karen and T. Gift. 2015. “Does Politics Influence Hiring? Evidence from a Randomized Experiment.” *Political Behavior* 37:653–675.
- Govtrack. 2021. “2020 Report Cards.”
URL: <https://www.govtrack.us/congress/members/report-cards/2020>
- Grimmer, Justin and Brandon Stewart. 2013. “Text as Data: The Promise and Pitfalls of Automatic Content Analysis Methods for Political Texts.” *Political Analysis* 21(3):267–297.
- Harbridge, Laurel. 2015. *Is Bipartisanship Dead? Policy Agreement and Agenda-Setting in the House of Representatives*. New York: Cambridge University Press.
- Hopkins, Daniel and Hans Noel. 2022. “Trump and the Shifting Meaning of “Conservative”: Using Activists’ Pairwise Comparisons to Measure Politicians’ Perceived Ideologies.” *American Political Science Review* pp. 1–8.
- Hutto, C.J. and Eric Gilbert. 2014. “VADER: A Parsimonious Rule-based Model for Sentiment Analysis of Social Media Text.” *Eighth International Conference on Weblogs and Social Media* .
- Itkowitz, Colby and Mike DeBonis. 2021. “Hostility Between Congressional Democrats and Republicans Reaches New Lows Amid Growing Fears of Violence.”
- Iyengar, Shanto, Guarav Sood and Yphtach Lelkes. 2012. “Affect, Not Ideology: A Social Identity Perspective on Polarization.” *Public Opinion Quarterly* 76(3):405–431.
- Iyengar, Shanto and Sean Westwood. 2015. “Fear and Loathing Across Party Lines.” *American Journal of Political Science* 59(3):690–707.
- Jurafsky, Dan and James Martin. 2021. *Speech and Language Processing: An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition*. Pearson Prentice Hall.
- Kaplan, Rebecca. 2021. “Nancy Pelosi: ‘The Enemy Is Within’ the House of Representatives.”
URL: <https://www.cbsnews.com/news/nancy-pelosi-enemy-within-house-of-representatives/>
- Kousser, Thad and Justin Phillips. 2012. *The Power of American Governors: Winning on Budgets and Losing on Policy*. Cambridge University Press.
- Le, Quoc and Tomas Mikolov. 2014. “Distributed Representations of Sentences and Documents.” *Proceedings of the 31st International Conference on Machine Learning* 32(2):1188–1196.

- Lee, Frances. 2016. *Insecure Majorities: Congress and the Perpetual Campaign*. University of Chicago Press.
- Levendusky, Matthew. 2009. *The Partisan Sort: How Liberals Became Democrats and Conservatives Became Republicans*. University of Chicago Press.
- Lewis, Jeffrey, Keith Poole, Howard Rosenthal, Adam Boche, Aaron Rudkin and Luke Sonnet. 2021. “Voteview: Congressional Roll-Call Votes Database.”
URL: <https://voteview.com>
- Mayhew, David. 1974. *Congress: The Electoral Connection*. Yale University Press.
- McConnell, Christopher, Yotam Margalit, Neil Malhotra and Matthew Levendusky. 2018. “The Economic Consequences of Partisanship in a Polarized Era.” *American Journal of Political Science* 62(1):5–18.
- Mutz, Diana and Byron Reeves. 2005. “The New Videomalaise: Effects of Televised Incivility on Political Trust.” *American Political Science Review* 99(1):1–15.
- Pedregosa, F., G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos, D. Cournapeau, M. Brucher, M. Perrot and E. Duchesnay. 2011. “Scikit-learn: Machine Learning in Python.” *Journal of Machine Learning Research* 12:2825–2830.
- Poole, Keith and Howard Rosenthal. 1997. *Congress: A Political-Economic History of Roll Call Voting*. New York: Oxford University Press.
- Quinn, Melissa. 2021. “House Votes to Censure Congressman Paul Gosar for Violent Video in Rare Formal Rebuke.” *CBS News* .
URL: <https://www.cbsnews.com/news/paul-gosar-censure-committees-house-aoc-video/>
- Řehůřek, Radim and Petr Sojka. 2010. Software Framework for Topic Modelling with Large Corpora. In *Proceedings of the LREC 2010 Workshop on New Challenges for NLP Frameworks*. Valletta, Malta: ELRA pp. 45–50. <http://is.muni.cz/publication/884893/en>.
- Rennie, Jason, Lawrence Shih, Jaime Teevan and David Karger. 2003. “Tackling the Poor Assumptions of Naive Bayes Classifiers.” *Proceedings of the Twentieth International Conference on Machine Learning* .
- Rohde, David. 1991. *Parties and Leaders in the Postreform House*. University of Chicago Press.
- Ruwitch, John and Barbara Sprunt. 2021. “Wyoming GOP Censures Liz Cheney for Voting to Impeach Trump.”
URL: <https://www.npr.org/2021/02/06/964933035/wyoming-gop-censures-liz-cheney-for-voting-to-impeach-trump>

- Sarlin, Benjy. 2021. "Some Democrats in Congress are Worried Their Colleagues Might Kill Them." *NBC News* .
URL: <https://www.nbcnews.com/politics/congress/some-democrats-congress-are-worried-their-colleagues-might-kill-them-n1254319>
- Slotkin, Jason. 2021. "South Carolina GOP Censures Tom Rice Over Trump Impeachment Vote."
- Sprunt, Barbara. 2021. "GOP Ousts Cheney From Leadership Over Her Criticism of Trump."
URL: <https://www.npr.org/2021/05/12/995072539/gop-poised-to-oust-cheney-from-leadership-over-her-criticism-of-trump>
- Stokes, Donald. 1963. "Spatial Models of Party Competition." *American Political Science Review* 57(2):368–377.
- The Lugar Center and the McCourt School of Public Policy. 2021. "Bipartisan Index."
URL: <https://www.thelugarcenter.org/ourwork-Bipartisan-Index.html>
- Volden, Craig and Alan Wiseman. 2018. "Legislative Effectiveness in the United States Senate." *Journal of Politics* 80(2).
- Zanona, Melanie. 2019. "Republican Bomb-Throwers Prep Impeachment Spectacle." *Politico* .
URL: <https://www.politico.com/news/2019/12/03/republican-lawmakers-impeachment-defend-trump-074988>

A Partisan Tweet Regex Dictionaries

In [Table 1](#), I have included the final regex dictionary used to identify partisan tweets. For Senator Mitch McConnell’s account, I excluded the terms that referred to himself as his account was a press account, which referred to him to reference his personal accomplishments for the state of Kentucky rather than to discuss him as a partisan referent. Additionally, for tweets mentioning the Republican President, I excluded tweets that mentioned universities, unless they explicitly mentioned Trump, POTUS, or the White House because there were many tweets which referred to university presidents.

Table A.1. Dictionary Identifying Partisan Tweets

Republican Party	Republican President	Republican Leadership	Democratic Party	Democratic President	Democratic Leadership
republican	trump	mcconnell	(?!anti)(?!un)democrat(?!ic government)(?!ic rights)(?!ic nation)(?!ic ally)(?!ic allies)(?!ic reform)(?!ic system)(?!ic ideal)(?!ic value)(?!ic norms)(?!ic process)(?!ic election)(?!ic institution)(?!ically)(?!ic participation)(?!ic freedom)(?!ic society)(?!ic principle)(?!ic republic)	biden	pelosi
gop(?!leader)(?!her)(?!arks)(?!anther)(?!olice)(?!ackgo)(?!ats)	((?!new)(?!next)(?!previous)(?!last)(?!vice)(?!senate)presiden(?!ts day)(?!tsday)(?!tial)(?!t elect)(?!t biden)(?:t of the (united states us)) (?!t of))(?!t and ceo)(?!t franklin)(?!t roosevelt)(?!t fdr)(?!t lbj)(?!t lyndon)(?!t johnson)(?!t harry)(?!t theodore)(?!t coolidge)(?!t calvin)(?!t george)(?!t john)(?!t james)(?!t madison)(?!t kennedy)(?!t jfk)(?!t abraham)(?!t lincoln)(?!t obama)(?!t @barackobama)(?!t truman)(?!t dwight)(?!t woodrow)(?!t eisenhower)(?!t wilson)(?!t reagan)(?!t ronald)(?!t carter)(?!t jimmy)(?!t ford)(?!t gerald)(?!t nixon)(?!t richard)(?!t xi)(?!t pro tem)(?!t bush)(?!t barack)(?!t bill)(?!t clinton)(?!t truman)(?!t washington)(?!t ulysses)(?!t grant)(?!t @rterdogan)(?!t erdogan)(?!t bolsonaro)(?!t @jairbolsonaro)(?!t jair)(?!t zelensky)(?!t guaido)(?!t @jguaido)(?!t shinzo)(?!t abe))	(sen senate)maj (ldr leader)	dems	obama	schumer
rnc	wh	gopleader	(?!tan)(?!epi)dem(\$[^\^a-z])	(?!mc)clinton(county co)	bidenharris
	(white house(?!hold))	(?!rs)(?!m)vp (?!biden)	(?!covi)dnc	barack	(biden harris)
	(?!sen)(?!senator)(?!sen)whitehouse	vpotus		((previous last new next) admin(stration))	(senkamalaharris)
	((?!new)(?!next)(?!previous)(?!last)(?!my)(?!our)(?!biden)(?!obama)(?!biden harris)(?!business)(?!insurance)(?!security)(?!veterans)(?!development)(?!drug)(?!services)admin(s stration))	(?!s)pence(?!r)		(president elect)	(senator harris)
	potus	mccarthy(?!ism)		((previous last new next)presiden (?!t and ceo)(?:t of the (united states us)))(?!t of))	kamala
	(commander in chief)	mitch(\$[^\^a-z])			((next new) ((vice president) vp))
	((^\)pre(s z)(?! obama)(?! biden)(?! @barackobama))	((?!next)(?!new)vice president(?! biden)(?! elect))			((((vice president elect) vp elect vp) harris)

B Test-Set Model Metrics

Table B.1. Democratic Test-Set Validation Metrics

Vectorization	Model	Accuracy	Balanced Accuracy	Prec. Neg.	Prec. Neut.	Prec. Pos.	Rec. Neg.	Rec. Neut.	Rec. Pos.	F1 Neg.	F1 Neut.	F1 Pos.
Doc2Vec	RandomForest	0.758	0.544	0.772	0	0.744	0.819	0	0.814	0.795	0	0.777
TFIDF	MultinomialNB	0.805	0.579	0.854	0	0.763	0.822	0	0.914	0.838	0	0.832
Uni+Bi-TFIDF	MultinomialNB	0.81	0.583	0.889	0	0.749	0.805	0	0.943	0.845	0	0.835
Doc2Vec	Perceptron	0.783	0.587	0.793	0.364	0.787	0.849	0.091	0.821	0.82	0.145	0.804
Doc2Vec	SVC	0.813	0.589	0.818	1	0.807	0.893	0.023	0.853	0.854	0.044	0.829
BOW	SVC	0.813	0.59	0.831	1	0.794	0.876	0.023	0.871	0.853	0.044	0.831
Uni+Bi-BOW	SVC	0.813	0.596	0.823	0.667	0.804	0.889	0.045	0.853	0.855	0.085	0.828
Uni+Bi-TFIDF	ComplementNB	0.821	0.597	0.866	1	0.781	0.846	0.023	0.921	0.856	0.044	0.845
BOW	RandomForest	0.823	0.597	0.847	1	0.798	0.876	0.023	0.892	0.861	0.044	0.843
Uni+Bi-BOW	RandomForest	0.815	0.598	0.85	0.5	0.785	0.856	0.045	0.892	0.853	0.083	0.836
Uni+Bi-TFIDF	Multinomial	0.824	0.598	0.838	0.5	0.813	0.886	0.023	0.885	0.861	0.043	0.847
Uni+Bi-BOW	MultinomialNB	0.824	0.599	0.861	1	0.79	0.856	0.023	0.918	0.859	0.044	0.849
BOW+TFIDF	RandomForest	0.816	0.599	0.843	0.667	0.792	0.862	0.045	0.889	0.852	0.085	0.838
Uni+Bi-TFIDF	SVC	0.829	0.601	0.839	1	0.818	0.893	0.023	0.889	0.865	0.044	0.852
Uni+Bi-TFIDF	RandomForest	0.815	0.604	0.839	1	0.79	0.856	0.068	0.889	0.847	0.128	0.836
TFIDF	SVC	0.839	0.608	0.84	1	0.838	0.913	0.023	0.889	0.875	0.044	0.863
TFIDF	ComplementNB	0.813	0.616	0.834	0.714	0.795	0.859	0.114	0.875	0.846	0.196	0.833
TFIDF	Multinomial	0.841	0.616	0.846	1	0.833	0.906	0.045	0.896	0.875	0.087	0.864
BOW	ComplementNB	0.816	0.624	0.815	0.75	0.82	0.886	0.136	0.849	0.849	0.231	0.835
Uni+Bi-BOW	ComplementNB	0.829	0.627	0.855	0.714	0.806	0.872	0.114	0.896	0.864	0.196	0.849
BOW	MultinomialNB	0.826	0.631	0.845	0.75	0.809	0.876	0.136	0.882	0.86	0.231	0.844
Uni+Bi-TFIDF	Perceptron	0.831	0.647	0.842	0.533	0.834	0.896	0.182	0.864	0.868	0.271	0.849
Doc2Vec	Multinomial	0.815	0.648	0.826	0.4	0.84	0.889	0.227	0.828	0.856	0.29	0.834
TFIDF	SVC(SGD)	0.821	0.659	0.821	0.458	0.853	0.893	0.25	0.835	0.855	0.324	0.844
BOW	SVC(SGD)	0.805	0.661	0.831	0.406	0.822	0.859	0.295	0.828	0.845	0.342	0.825
TFIDF	Perceptron	0.824	0.662	0.832	0.55	0.835	0.883	0.25	0.853	0.857	0.344	0.844
Uni+Bi-BOW	SVC(SGD)	0.833	0.674	0.861	0.522	0.827	0.876	0.273	0.875	0.869	0.358	0.85
Uni+Bi-BOW	Perceptron	0.818	0.676	0.827	0.583	0.828	0.883	0.318	0.828	0.854	0.412	0.828
Doc2Vec	SVC(SGD)	0.792	0.677	0.834	0.333	0.832	0.846	0.386	0.799	0.84	0.358	0.815
Uni+Bi-BOW	Multinomial	0.829	0.678	0.853	0.52	0.83	0.896	0.295	0.842	0.874	0.377	0.836
Uni+Bi-TFIDF	SVC(SGD)	0.839	0.685	0.845	0.5	0.864	0.913	0.295	0.846	0.877	0.371	0.855
BOW	Multinomial	0.831	0.685	0.849	0.452	0.853	0.903	0.318	0.835	0.875	0.373	0.844
BOW	Perceptron	0.813	0.686	0.841	0.471	0.824	0.872	0.364	0.821	0.857	0.41	0.822

Table B.2. Republican Test-Set Validation Metrics

Vectorization	Model	Accuracy	Balanced Accuracy	Prec. Neg.	Prec. Neut.	Prec. Pos.	Rec. Neg.	Rec. Neut.	Rec. Pos.	F1 Neg.	F1 Neut.	F1 Pos.
Uni+Bi-TFIDF	MultinomialNB	0.66	0.364	0.653	1	1	1	0.011	0.08	0.79	0.022	0.149
TFIDF	MultinomialNB	0.694	0.409	0.683	0	0.843	0.997	0	0.229	0.811	0	0.36
Doc2Vec	RandomForest	0.708	0.432	0.704	1	0.739	0.993	0.011	0.291	0.824	0.022	0.418
Uni+Bi-TFIDF	ComplementNB	0.716	0.445	0.701	1	0.861	0.995	0.028	0.312	0.822	0.054	0.458
Uni+Bi-BOW	MultinomialNB	0.771	0.525	0.761	0.75	0.819	0.984	0.033	0.558	0.859	0.063	0.664
Uni+Bi-BOW	RandomForest	0.775	0.54	0.767	1	0.806	0.979	0.066	0.573	0.86	0.124	0.67
Uni+Bi-TFIDF	RandomForest	0.78	0.542	0.772	1	0.811	0.979	0.044	0.603	0.864	0.085	0.692
TFIDF	RandomForest	0.775	0.543	0.771	1	0.787	0.967	0.05	0.613	0.858	0.095	0.689
BOW	RandomForest	0.783	0.557	0.792	1	0.746	0.959	0.05	0.663	0.867	0.095	0.702
TFIDF	ComplementNB	0.777	0.558	0.787	0.677	0.747	0.966	0.116	0.593	0.868	0.198	0.661
Uni+Bi-TFIDF	SVC-OVR	0.808	0.587	0.809	0.882	0.802	0.978	0.083	0.701	0.885	0.152	0.748
Doc2Vec	SVC-OVR	0.798	0.591	0.826	0.808	0.712	0.956	0.116	0.701	0.886	0.203	0.706
BOW	SVC-OVR	0.797	0.594	0.813	0.885	0.74	0.953	0.127	0.701	0.878	0.222	0.72
Uni+Bi-BOW	ComplementNB	0.805	0.6	0.805	0.756	0.809	0.977	0.171	0.651	0.883	0.279	0.721
Doc2Vec	SVC(SGD)	0.788	0.604	0.827	0.576	0.704	0.938	0.21	0.663	0.879	0.308	0.683
Doc2Vec	Perceptron	0.743	0.606	0.844	0.384	0.62	0.858	0.32	0.638	0.851	0.349	0.629
Uni+Bi-TFIDF	Multinomial	0.815	0.607	0.822	0.852	0.786	0.973	0.127	0.721	0.891	0.221	0.752
TFIDF	SVC-OVR	0.813	0.608	0.83	0.846	0.759	0.965	0.122	0.736	0.892	0.213	0.747
TFIDF	Multinomial	0.822	0.64	0.85	0.804	0.744	0.955	0.204	0.761	0.9	0.326	0.753
BOW	MultinomialNB	0.824	0.655	0.858	0.712	0.746	0.951	0.26	0.754	0.902	0.381	0.75
BOW	SVC(SGD)	0.808	0.658	0.867	0.535	0.722	0.933	0.337	0.704	0.899	0.414	0.712
TFIDF	SVC(SGD)	0.825	0.658	0.853	0.722	0.76	0.956	0.287	0.731	0.901	0.411	0.745
Doc2Vec	Multinomial	0.792	0.66	0.856	0.55	0.694	0.904	0.392	0.683	0.879	0.458	0.689
Uni+Bi-TFIDF	Perceptron	0.82	0.666	0.92	0.85	0.642	0.895	0.188	0.915	0.907	0.308	0.754
BOW	Perceptron	0.795	0.672	0.888	0.451	0.696	0.895	0.409	0.714	0.891	0.429	0.705
BOW	ComplementNB	0.822	0.675	0.88	0.652	0.712	0.928	0.32	0.776	0.903	0.43	0.743
Uni+Bi-BOW	Perceptron	0.823	0.675	0.876	0.612	0.731	0.935	0.331	0.759	0.905	0.43	0.745
BOW	Multinomial	0.82	0.677	0.887	0.559	0.719	0.926	0.343	0.764	0.906	0.425	0.741
TFIDF	Perceptron	0.81	0.679	0.883	0.544	0.707	0.91	0.376	0.751	0.897	0.444	0.728
Uni+Bi-BOW	Multinomial	0.838	0.683	0.881	0.747	0.743	0.952	0.309	0.786	0.915	0.438	0.764
Uni+Bi-BOW	SVC(SGD)	0.835	0.685	0.88	0.667	0.75	0.946	0.331	0.776	0.912	0.443	0.763
Uni+Bi-TFIDF	SVC(SGD)	0.848	0.697	0.886	0.803	0.759	0.952	0.315	0.824	0.918	0.452	0.79

C Validating the Model Predictions

C.1 Random Samples of 10 Tweets from Each Predicted Category

This section includes random samples of 10 tweets predicted to be in each category. As we can see, the tweets themselves largely make sense and fit with their categories. The first tweet from Anthony Brown, for example, calls Biden’s pick for Secretary of Defense “historic and significant.”

Table C.1.1. Random Sample of 10 Predicted Positive Democratic Tweets

Username	Tweet
RepAnthonyBrown	Historic and significant - if he is @JoeBiden's pick for SecDef Gen Lloyd Austin has the character and competence necessary to lead the Department of Defense\n\nLloyd Austin is top flight and he's the right choice to lead our civilian & military personnel at the Pentagon https://t.co/h2DQkpdmNV
RepLoisFrankel	We have to do more for the countless victims of #gunviolence in our nation. This is a crisis - Trump, when will you and Republicans join with Democrats to finally pass the commonsense gun safety legislation we so desperately need? #SOTU
SenMarkey	Proud to stand with @SpeakerPelosi, @SenSchumer, my @SenateDems colleagues and other lawmakers today to announce the Save The Internet Act. Let's restore #NetNeutrality and return the internet to the American people. #SaveTheNet https://t.co/Q5QA21fLPY
RepJoeNeguse	As extreme weather events become more and more frequent, we must be prepared to mitigate and rebuild. This week, @HouseDemocrats are pushing a new disaster assistance package to add \$17.2 billion in relief & recovery assistance for Americans affected by recent natural disasters.
RepKClark	If the White House and @RussVought45 have nothing to hide, why won't they comply with the @HouseDemocrats process and get the facts before the American people? #TruthExposed https://t.co/cmCp4SQ4bC
repjimcooper	Congratulations @JoeBiden & @KamalaHarris. Let's get to work!
RepCartwright	The @AppropsDems Labor-HHS-Education bill we're considering today funds some of our nation's most critical programs. It includes robust funding for state & local public health departments for #COVID19 response, child care programs, and assistive services for older Americans. https://t.co/YdIxRJP8qq Watch here https://t.co/njaLVHhEBB
RepAngieCraig	As a member of @TransportDems, I've been advocating for important projects in our community. I'm grateful that the INVEST in America Act includes potential key #MN02 transportation priorities. https://t.co/BcOy5WWemc
RepSmucker	Today, I spoke against House Democrats' efforts to roll back the student loan borrower defense repayment regulation, offered by President Trump's Administration. Thankfully, House Democrats' efforts were unsuccessful. Watch my comments below https://t.co/q0s32p7pFB
RepJimmyGomez	More money in #WorkingFamilies' pockets has ALWAYS been a priority for @HouseDemocrats. Now that we've passed this bill, it's time for @senatemajldr to choose: Join Democrats & send this legislation to @realDonaldTrump's desk or block another relief bill for struggling Americans. https://t.co/gBoZAF5IfN Here is the official vote for the CASH ACT. I voted Yea!\n\n https://t.co/fJhpyz5xxm

Table C.1.2. Random Sample of 10 Predicted Neutral Democratic or Non-Democratic Tweets

Username	Tweet
RepBrindisi	ICYMI: We worked with Democrats and Republicans to get a trade deal that works for hardworking Upstate New York farmers and businesses. Looking forward to voting on final passage of the USMCA soon. https://t.co/n0Cfcu7l9
SenBobCasey	"The American people should have a voice in the selection of their next Supreme Court Justice. Therefore, this vacancy should not be filled until we have a new president." \n \n- Senate Majority Leader McConnell, February 2016.
RepValDemings	.@SpeakerPelosi is right: the president and his associates are engaged in a cover up.
RepRoKhanna	To truly make a change in American policing, any legislative solutions out of Congress need to be bipartisan. @RepGallagher laid out 2 areas where Republicans and Dems might be able to come together: \n \n1. De-militarizing the police \n \n2. Ending qualified immunity https://t.co/nNPLmCzA8T
GovNedLamont	On the set with the gang of #CapitolReport, including this young man's mother, @mrskurantowicz. She's a Republican, and I'm a Democrat, but her son Iggy is just plain cool. @ctcapitolreport @WTNH https://t.co/ThUihTFaZN
RodneyDavis	He's a Democrat, I'm a Republican. I invited him to DC to share his unique perspective as the clerk of a small county who administers elections because election security shouldn't be a partisan issue. Thank you Christian Co Clerk Mike Gianasi your insight https://t.co/RUpZWBApQr
RepMattGaetz	"Democrats and Republicans may not agree on a wide variety of issues, but as congressmen from both parties, we agree with Ambassador Espina on one: We need to contain the influence of China in the Western Hemisphere." \n \nvia @RepGonzalez, Amb. Espina, and I: \n https://t.co/JKhxFUQkk6
RepBrindisi	Beautiful Memorial Day morning in Clinton honoring those who gave the ultimate sacrifice in service to our county. https://t.co/NK9Evs624C
RandPaul	What brings Big Government Republicans and Democrats together? Support for Endless War. After 19 years in Afghanistan, it's high time to bring our troops home! \n \n https://t.co/I14DOYACct
RepFredUpton	Folks really do want us to work together - Republicans & Democrats. @RepDebDingell & I are two vice chairs of the @ProbSolveCaucus, and we've taken the lead on a number of issues as we really try to get things done. That's what people want us to do. #MI06 https://t.co/TFXOM2nMJP

Table C.1.3. Random Sample of 10 Predicted Negative Democratic Tweets

Username	Tweet
RepAndyBiggsAZ	COMING UP: I'm joining @JonScottFNC on @AmericaNewsroom to discuss the latest on the Democrats' partisan impeachment plot + how Dems continue their efforts to block @POTUS @realDonaldTrump's attempts to solve the humanitarian crisis at the border. Watch @FoxNews at 9:10 EST.
RonWyden	While Democrats push for aggressive action against Russian assets interfering in our elections Republicans are parroting them to advance bogus investigations. Senate investigations shouldn't rely on conspiracy theories pushed by shady characters trying to undermine our democracy. Andrii Derkach has been central in advancing the Russian disinformation that underpins Senate Republicans' effort to smear Vice President Biden. He is now under U.S. sanctions for his efforts to interfere in the election. https://t.co/xQZIH9ZMVI
GReschenthaler	Nancy Pelosi and Chuck Schumer spent the last 4 months prioritizing stimulus checks for illegal immigrants over targeted relief for struggling Americans.\n\nAmericans see right through their games and they're fed up. That's why you saw so many House Democrats lose their jobs. https://t.co/AlX4x0qqG5
RepFredKeller	If you do not believe the @HouseDemocrats are putting partisan politics first, think of this ridiculous contrast:\n\nYesterday, they blocked a vote to support the freedom-seeking Iranian protesters.\n\nToday, they will continue their sham impeachment of President @realDonaldTrump. https://t.co/rdEnliSWlu
RepFredKeller	Impeachment is not a tool for exacting political vengeance.\n\nSpeaker Pelosi's indication that she is willing to revive the Democrats' failed impeachment sham is nothing but a cheap attempt to obstruct President Trump from fulfilling his Constitutional responsibility. Disgraceful.
CongMikeSimpson	Science should be transparent and I don't understand how Democrats could oppose a rule that is designed to make @EPA's science publicly available so we can better understand how & why they create rules & regulations. \n\nWatch my comments to learn more. https://t.co/fgJ0ihQ3hx
RepChuck	15,000 good-paying jobs that support hardworking East Tennessee families. What could make these jobs even better?\n\nPassing the #USMCA to provide long-term certainty to the folks who deserve it the most. @SpeakerPelosi your call. https://t.co/kcgrcqZToC
FrankPallone	I joined 23 other @EnergyCommerce Democrats in a letter to @Ajit-PaiFCC urging the @FCC to delay a vote on a Declaratory Ruling that would limit local governments' role in the deployment of wireless infrastructure.\n https://t.co/drON30JsKZ
RepGosar	Democrats care more about tearing @realDonaldTrump down than building America up! https://t.co/mSgqAPRxgd
SenRickScott	Florida and Puerto Rico are still waiting on important disaster relief funding. Tonight we'll vote on the bill I'm co-sponsoring.\n\nI urge Democrats to put politics aside and vote for this important legislation. American families need this funding NOW!

Table C.1.4. Random Sample of 10 Predicted Positive Republican Tweets

Username	Tweet
RepAndyBiggsAZ	After months of economic nightmares due to COVID-19 & states' reactions to the outbreak, our economy is roaring back as predicted. \n\nAmericans trust President @realDonaldTrump & his team to lead our economy back to prosperity because of his track record. https://t.co/kBG2iwWOb0 @realDonaldTrump The fundamentals of our economy were very strong prior to the COVID-19 outbreak, and the pro-growth foundation that President Donald J. Trump set over the past 3+ years is paying dividends in one of our nation's most-uncertain times. @realDonaldTrump As our economy is restored, it is imperative that President Trump is not undermined in his mission to return our economy to greatness. @realDonaldTrump Dr. Anthony Fauci and Dr. Deborah Birx continue to contradict many of President Trump's stated goals and actions for returning to normalcy as we know more about the COVID-19 outbreak. This is causing panic that compromises our economic recovery. @realDonaldTrump We can protect our most vulnerable from COVID-19 while still protecting lives & livelihoods of the rest of the population. \n\nIt's time for the COVID-19 task force to be disbanded so that President Trump's message is not mitigated or distorted. https://t.co/gG3oIj7KI0
SenatorBraun	As Democrats descend on Miami for their first presidential #debate, it's a good time to reflect on the fact that President @realDonaldTrump has kept his promise to nominate great conservative judges, with the Senate approving 125 since Trump took office.\n https://t.co/a70QzcP6OM
DesJarlaisTN04	Hillary funneled foreign donations through State Department to Clinton Foundation. Destroyed evidence. Quid Pro Joe Jr. profited in Ukraine while VP Biden distributed foreign aid (after failed 'Russia Reset'). Donald Trump promises to Drain the Swamp – Democrats attack!
JacksonLeeTX18	...which included only democratic votes for the stimulus. That's why we have the progress we see today - democrats working with President Obama. One of the challenges of the 164k jobs report is that these jobs are highly technical ones. This morning's report detailing 164k jobs is continuing the trajectory of the Obama administration where Democratic members of congress worked with the administration to get the United States out of the financial abyss left by the previous president... The question is: what do we do to address automation and teaching our young children to be trained and ready for 21st century jobs. The administration does not know. Democrats must act.
RepHarley	My provisions, which increase competition and overhaul outdated regulations, would improve Southern California's transportation and energy systems while saving taxpayer dollars. Read more (2/x) https://t.co/m0zLA1bhhp 25 hours later, @TransportDems and @TransportGOP have passed the INVEST in America Act, which includes four of my provisions. This legislation is a strategic and cost-effective bill that empowers businesses, protects our environment, and creates quality American jobs. (1/x) Throughout the coronavirus crisis, we have been working #ForThePeople. Even via WebEx, I will continue fighting for working families and small businesses across Orange County. (3/3)
RepCartwright	Children's hospitals have been greatly challenged by #COVID19, but they're getting seriously shortchanged when it comes to relief funding. PA Democrats & Republicans agree that needs to change, and fast. Thanks to @RepSusanWild & @RepBrianFitz for leading this bipartisan effort. https://t.co/yQgaDkvGr7
RepSteveChabot	WATCH @stevenmnuchin1 and @SBAJovita testify in front of @HSBCgop on the status of small businesses impacted by #COVID19 #PPP #EIDL here: https://t.co/DNHnzApj1r
RepMarkTakano	Amb. Sondland donated \$1 million to Trump's inaugural committee and got a cushy job as Ambassador to the EU.\n\nHe's a staunch supporter of the president and admitted that there was a "quid pro quo" to force Ukraine to launch an investigation into Trump's rival.\n\nIn his own words https://t.co/RylwxxGYnD
RepRickCrawford	The President's actions today forming public-private partnerships in our fight against the #COVID virus will bring us victory, much like we experienced during WWII. The United States of America is at its best when we are united in purpose. #America #United\n https://t.co/g4AmYn6Oof
GOPLLeader	Happy birthday to President Trump! https://t.co/xhzB8yxQCo

Table C.1.5. Random Sample of 10 Predicted Neutral Republican or Non-Republican Tweets

Username	Tweet
RepBeatty	I call on @realDonaldTrump to approve the State of Ohio's Major Disaster Declaration ASAP! #COVID19 https://t.co/nHE2NY9pti
RepHarley	Today, I joined @newportchamber of Commerce President, Steve Rosansky, to hear directly from local small business owners who are struggling during #COVID19. Small businesses are the backbone of our nation – we must ensure they can survive the coronavirus crisis. https://t.co/w6yjwtC4w1W
ChrisCoons	I remember four years ago today, my colleagues - Democrats and Republicans - came together to thank @JoeBiden for decades of bipartisan leadership and public service. \n\nLet's all remember how we've worked with Joe Biden and we can again. \n\nThat's how we make people's lives better. https://t.co/hTIhU3k9rD
Sen_JoeManchin	My bill, the Great American Outdoors Act has passed in the House & is now headed to the President's desk! This bipartisan, landmark conservation legislation will protect & invest in our nation's public lands. I look forward to @realDonaldTrump signing this legislation into law. https://t.co/R7Gx8xAkUO
RepRoybalAllard	ALERT: Are you coming to DC at #Easter time and want to attend the White House Easter Egg Roll on April 13? Tickets are FREE to the public through an online lottery at https://t.co/nUMvt7ouFl . Be sure to enter by Monday, February 24 at 7AM PT / 10AM ET!
SanfordBishop	It's disheartening to hear that the expected disaster relief, that was on its way to @POTUS' desk, has been delayed yet again. However, I am hopeful that the House will pass this vital piece of legislation next week!
SenSherrodBrown	WATCH LIVE: Joining @SenBobCasey @RonWyden @SSWorks @CAPDisability and @LittleLobbyists to tell President Trump: #NoSocialSecurityCuts https://t.co/8uoEecWrwl
RepRaskin	"Swastikas and Confederate flags, nooses and automatic rifles do not represent who we are..." Solidarity with Governor Whitmer and Michiganders facing down rabid gun-toting Trump-inflamed zealots. https://t.co/MH1fayoPBU
dougducey	Congratulations to all the @NAU teachers academy graduates! They'll be bringing the skills they've learned to the front of #AZ classrooms #ThingsThatMatterAZ https://t.co/HrF7YjFAtw @AZGovEducation @NAU Thank you @NAUPresident Rita Cheng for helping make this program a success! @NAU
RepFredKeller	Back in February, the House passed my first bill: H.R. 4279 to name the Post Office in Laceyville, PA, as the "Melinda Gene Piccotti Office." \n\nI'm proud to announce that H.R. 4279 recently passed the Senate and is now headed to the President's desk to become law. https://t.co/ThlB32TP6X The late Mindy Piccotti founded Hunts for Healing to help veterans returning from combat. This bill celebrates Mindy's life and her continued legacy in our veteran community. \n\nMore on H.R. 4279: https://t.co/UL5Md90ukB

Table C.1.6. Random Sample of 10 Predicted Negative Republican Tweets

Username	Tweet
RepBarbaraLee	"I can't afford to go to the ER. I can't afford anything. I just went to bed and hoped I'd wake up." - Mallory Lorge, federal worker.\n \nLives are at risk in this irresponsible shutdown. Pres Trump needs to stop using families as pawns! #EndTheShutdown\nhttps://t.co/149eqFPPeK
SenJeffMerkley	Reminder: last time Donald Trump claimed he was going to get tough on drug companies, their stocks went UP. https://t.co/G624MbhFtB
RepBarbaraLee	We're coming back to Congress to stop Trump in his tracks & save the @USPS. \n \nOur democracy & security is at stake.
NormaJTorres	It's no surprise that @realDonaldTrump is more concerned w/ corporate profits than people's health.\n \nMr. President: It's your job to keep America safe. Gutting the @CDCgov budget, whose main job is to respond to threats like the #Coronavirus, is not how you do it.\n \nDo your job. https://t.co/2EYmLYp04Q
RepSpeier	This is a shameful example of a debate. It's not a debate. It's a childish slug fest incited by POTUS. Cut the mike when they interrupt or exceed time
RepLloydDoggett	Trump knew this virus was deadly early on, yet—even after contracting it—he continued spreading his lies and the virus itself at super-spreader events.\n \nI've tracked his ongoing denial, deception, and distractions here:\nhttps://t.co/ldMF0JoelW
BillPascrell	We've been demanding this action for days. It means we bring America's full might to produce whatever materials we need to defeat the virus. Now trump needs to activate America's full strength Today. https://t.co/kGLdepRs0C
SenSchumer	The @NYTimes reported that President @realDonaldTrump is ordering ICE to resume plans to carry out mass deportation raids over the weekend.\n \nHis plan will tear families apart and disrupt communities across America.\n \nCruelty seems to be the point of these #ICEraids. https://t.co/9NhZZkGEYJ
SenatorDurbin	We can't let Pres Trump & Republicans get away with their scheme to suppress the vote by dismantling @USPS.\n \nDemocrats will keep doing everything we can to make sure @USPS receives funding in the next #COVID19 relief bill so Americans can safely vote this fall. #DontMess-WithUSPS
RepMullin	Today, the Trump Administration is kicking off National Native American Heritage Month. This month gives all Americans the opportunity to celebrate the legacy of the first people who called this land home. Join the event here https://t.co/RCvorA1cv5

D Relationships Between Affect Scores and Other Variables

Figure D.1. House Democratic Correlation Plot

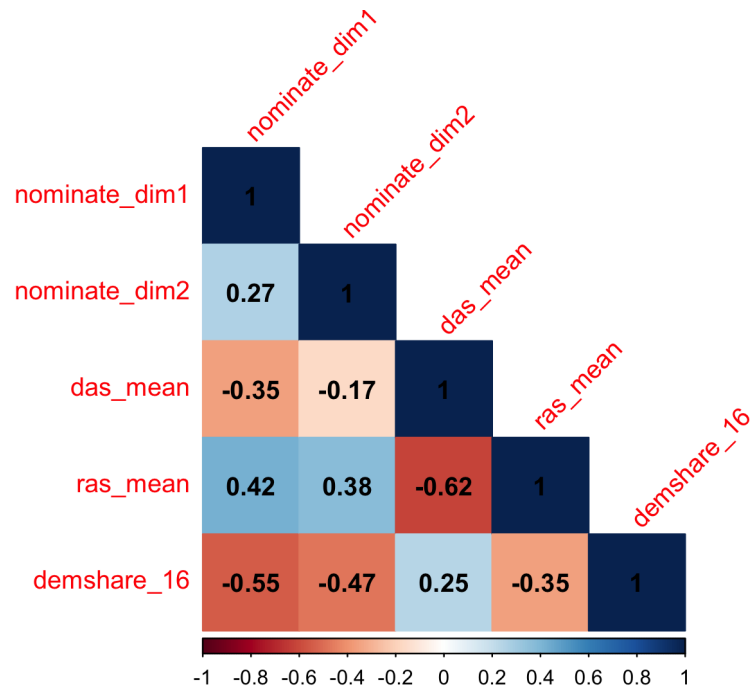


Figure D.2. House Republican Correlation Plot

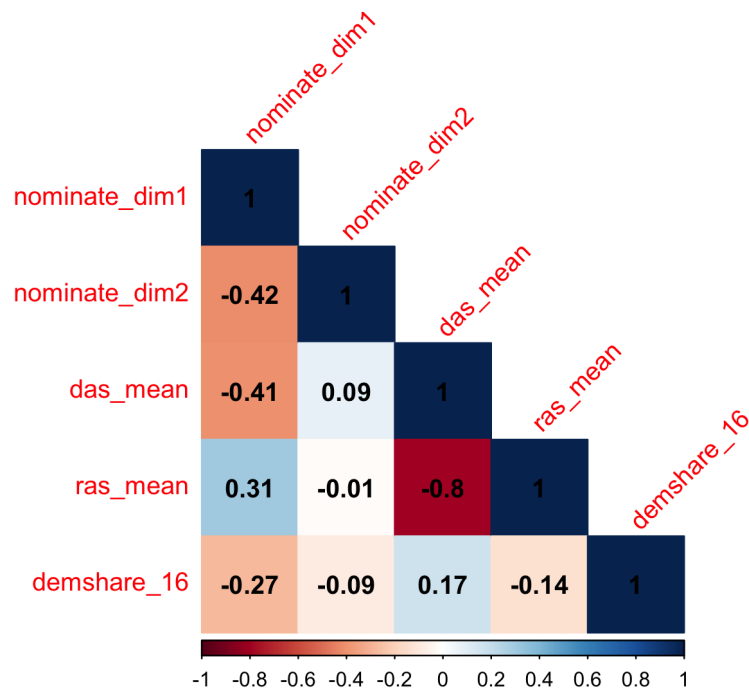


Figure D.3. Senate Democratic Correlation Plot

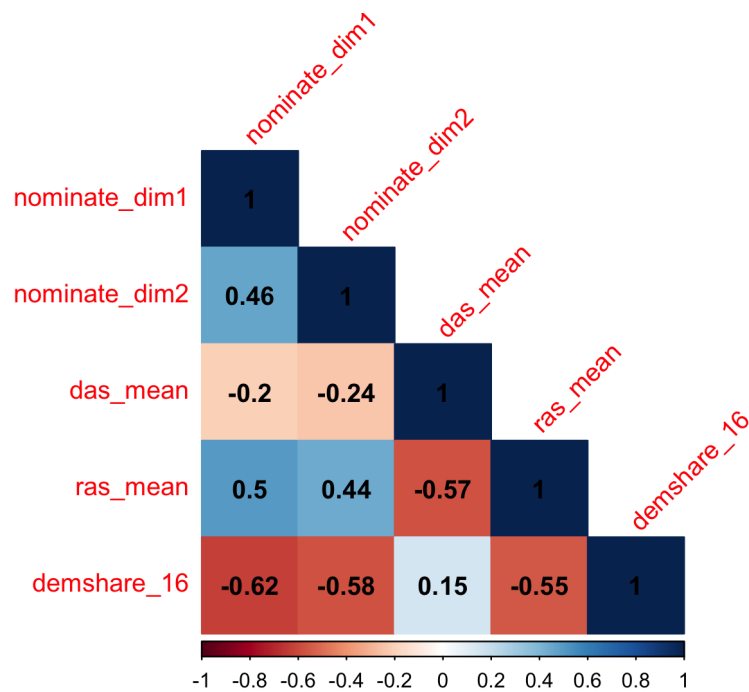


Figure D.4. Senate Republican Correlation Plot

