

Hooked With Phonetics: The Strategic Use of Style-Shifting in Political Rhetoric

Markus Neumann The Pennsylvania State University

The study of non-policy representation has emphasized written or verbal communication by representatives, neglecting the crucial non-verbal component of symbolic representation. I argue that, in order to convince their constituents that they are “like them” and will act in their interests, politicians project likeness through the way they talk. When speaking to an audience of ordinary citizens, a politician will play the “average Joe.” When addressing a congressional committee, they will demonstrate sophistication and competence. Hence, political speakers shift their style according to the audience. I test this hypothesis using audio data from congressional and campaign speeches of U.S. senators uploaded to YouTube. I extract the acoustic properties of the audio signal and measure vowel space density, a concept developed in phonetics, to categorize the degree of articulation in speech. The results suggest that politicians do adjust their articulation to fit the needs of their audience. This effect is greater for female senators, who also speak with a higher degree of formality in general.

Introduction

The representative-constituent relationship lies at the heart of every representative democracy: Citizens vote for candidates, who then turn their constituents’ political attitudes into policy. As a result, discourse, and even more so, quantitative scholarship, on representation is frequently limited to policy congruence and responsiveness: Does the policy that is being implemented match the preferences of the constituents? However, representation is much broader than this.

As noted by Pitkin (1967), representation is as much about “standing for” as it is about “acting for” a constituency. The former is harder to define, but nevertheless important and falls under the label of ‘symbolic representation’ (Eulau and Karps, 1977). As opposed to the transactional aspects of representation, symbolic representation is more about gestures than palpable actions (Hill and Hurley, 2002). Their goal is not necessarily to leave constituents with the impression that a representative acted in their interests, but more so that she is a “good person”. Symbolic representation occupies a critical role in some of the most important theories of political representation, be it in Congress (Mayhew, 1974) or at home (Fenno, 1978).

If the non-policy, non-transactional aspects of representation are so important, then why has scholarship in the period since Mayhew (1974), Eulau and Karps (1977) and Fenno

(1978) paid relatively little attention to it? Symbolic representation is a hard problem to study. As opposed to roll-call votes on policy, the nuances of social interactions between constituents and their representatives are difficult to conceptualize and measure. Grimmer (2010) – who attempts to further develop Fenno’s concept of home style – offers a solution by analyzing Senate press releases with the help of text as data methods.

However, in pursuit of the goal of obtaining a quantitative measure of home style, Grimmer adapts a limited version of Fenno’s theory. One of the important distinctions of Fenno (1978) is that he pays close attention to both verbal and non-verbal forms of communication. The latter, he argues, represents a more honest indicator of the true intentions of legislators and constituents therefore pay special attention to it. By contrast, Grimmer (2010) focuses only on the verbal. He dismisses the non-verbal forms of communication studied by Fenno as “folksy mannerisms”. However, I argue that these mannerisms are important and should not be overlooked. These acts of symbolic representation are aimed at sending the following message: “You can trust me because we are like one another” (Fenno, 1977).

In this research, I expand on this idea and develop a theory of likeness. I regard perceived likeness as a sufficient (but not necessary) condition for constituents to trust their representative. The electoral implication then, is implicit: “I am like you, therefore you should vote for me.” This is a departure from the classic concept of representation, where the message is more akin to: “I want the same things as you and therefore you should vote for me.” Likeness is a more diffuse concept than descriptive representation, which exists when representatives match a specific demographic qualifier of their constituents – i.e. African Americans politicians representing African Americans, women representing women, farmers representing farmers, and so on (Mansbridge, 1999). By contrast, likeness is about a *perceived* match in attributes between constituent and representative. Whether the two are objectively alike is not important, what matters is that the bond *feels* true. And as noted by a number of scholars (Mayhew, 1974; Fenno, 1978, 2007; Rogers and Nickerson, 2013; Hollibaugh, Rothenberg and Rulison, 2013; Grose, Malhotra and Parks Van Houweling, 2015), perception is absolutely critical for representation.

So how does a politician project likeness with her constituents? The strategy I focus on here is phonetic style-shifting. I argue that representatives rely on subtle rhetorical clues in order to signal likeness with their constituents. Specifically, I examine the modulation of *articulation* – the clarity and fullness with which phones, the fundamental building blocks of speech, are pronounced. By lowering their degree of articulation to a colloquial level, representatives signal to their voters that they are one of them. There are a number of rhetorical devices representatives can use to accomplish this. Regional dialects and accents are one approach. There are also a number of informal pronunciations, such as

“y’all” or g-dropping – the practice of omitting the /g/ from words ending in /-ing/. Sociolinguists have studied this particular phenomenon extensively and also observed politicians making use of it in certain contexts (Liberman, 2008; Nunberg, 2008; Liberman, 2011). I expect the direction of style-shifting to be dependent on the audience in front of which a politician performs: When politicians want to demonstrate competence and appeal to elites – as in, when they are speaking in Congress – they will use a higher degree of articulation. By contrast, when the primary goal is to demonstrate empathy and appeal to the “average Joe” – as would be the case in a campaign setting – representatives will use a lower degree of articulation. In my eyes, phonetic style-shifting thus represents a more faithful translation of the non-verbal component of Fenno’s concept of representation of self than Grimmer’s approach.

I rely on the U.S. Senate to test my theory. The upper chamber of the American legislature provides the ideal test case because senators represent diverse constituencies and command the national spotlight to an extent that provides sufficient data. In analyzing the audio of political speeches, I make use of a medium political science has largely ignored so far. My data is scraped from the YouTube channels of senators, most of whom maintain two channels - one for campaigns, and one for legislative activity. This division already suggests that my theory points in the right direction, and provides me with a convenient way of assessing the differences in speech patterns in low- (campaigns) versus high-brow (Congress) settings. My sample spans 64 senators who were part of the 115th or 116th Congress, and contains 3466 videos, which run for a combined 54 hours.

In keeping with prior research in phonetics (Sandoval et al., 2013; Story and Bunton, 2017), I rely on vowel space area, which approximates the extent to which the throat and mouth are used in the generation of sounds, to operationalize the concept of articulation. Lower articulation makes vowels less distinct from one another, thus compacting the vowel space. I compare the vowel space area of senators in two different settings - campaigns and Congress - and show that senators consistently make greater use of the vocal apparatus in the latter case. This difference is even greater for female senators, who also have a larger vowel space in general. This provides evidence in favor of my argument that a variation in articulation, contingent on the audience, allows politicians to perform acts of symbolic representation through the projection of likeness.

This research has important implications for the concept of representation. The study of non-policy representation, while not as common as policy representation, is well-established, and as such follows well-tread paths – either through the use of anecdotes, as with Fenno, or the examination of explicitly verbal forms of communication, as demonstrated by Grimmer (see also Maltzman and Sigelman (1996); Hill and Hurley (2002)). My research shows that the more subtle and nebulous aspects of non-policy

responsiveness can nevertheless be studied with quantitative evidence.

Furthermore, the strategy of phonetic style-shifting can provide some insight into the phenomenon of populism. Populism rails against the establishment, against the government and against the elite (Moffitt, 2016). To varying degrees, this is something all politicians engage in. As Fenno (1978) famously notes, “Members of Congress run for Congress by running against Congress.” But populism contains an inherent contradiction, as populists themselves are elites, and at the very least aim to form the government. Consequently they are faced with the challenge of convincing their voters – who hate elites – that they are one of them, and not part of the elite. My theory of likeness allows for politicians to perform both roles simultaneously and to shift between them fluidly by varying the degree of articulation in their speech.

Literature

The Perception of Likeness

Mayhew (1974) posits that in order to pursue the goal of getting reelected, politicians engage in three activities: (1) position-taking, (2) advertising and (3) credit-claiming. Ultimately, all of these strategies accomplish their goal by building a specific image of the politician in the eyes of his or her constituents. This image, reduced to its core, is: “I am like you and therefore you should vote for me.” In a way, this is a perversion (or in the words of Fenno (1977), “corruption”) of the classical view (APSA, 1950; Downs, 1957; Schattschneider, 1960) of the electoral connection, in which politicians (or their parties) essentially claim: “I want the same policies as you and therefore you should vote for me.” But as Mayhew (1974) observes, this does not actually happen. It is possible for members of Congress to vote against the interests of their constituents and get away with it because keeping track of what Congress does at all times and evaluating how its decisions fit into one’s preference system (to the extent that it even exists (Converse, 1964)) requires citizens to pay enormous information costs (see also Arnold (1990). Rather than supporting politicians who represent their policy interests, voters simply tend to adopt the policy positions of their representatives (Lenz, 2009; Broockman and Butler, 2017).

Politicians also exploit the fact that “I am like you” is a shortcut for “I want the same policies as you.” As noted by one of the politicians studied by Fenno (1977), as long as voters believe that their representative is “a nice fella,” they are willing to make “presumptions in my favor.” As a result, the way politicians present themselves to citizens is enormously important. Mayhew (1974) and Fenno (1977) have laid the groundwork here, as they have analyzed the lengths that representatives go to in order to cultivate

their image in the public eye – both when they spend their time in Washington, as well as in their districts. To that end, Fenno points to their presentation of self as being absolutely crucial. Fenno (1977) adapts this concept from the sociologist and social psychologist Erwin Goffman. Goffman (1959) points out that when it comes to communication, both verbal and non-verbal expressions are critical. In fact, between the two, Goffman attributes greater importance to non-verbal communication. Fenno (1977) translates this idea into the political realm: He points out that constituents cannot fully trust the performance of their representative: He may claim to be acting in their best interest, but his verbal assurances are no guarantee. As a result, they look towards his non-verbal behavior as a more honest indicator for his true intentions.

Since Fenno, several of scholars have attempted to expand on his concept of home style. For example, Yiannakis (1982) conducts a manual content analysis of newsletters and press releases under the criteria established by Mayhew (1974). McGraw, Timpone and Bruck (1993) focus on the response to justifications of members of Congress with respect to their voting choices. Lipinski (2004), elected to Congress shortly after the publication of his book, analyses district mailings and studies how party and majority membership influence the level of positivity representatives use in talking about the institution. However, such studies on political communication with the goal of extending the theory built by Fenno are far and few between. As pointed out by Grimmer (2013) – using a topic model on the literature citing Fenno – much of this work does not really engage his ideas (especially the style versus substance argument based on Goffman (1959)), and over time, has mostly shifted to institution-focused studies of Congress.

Grimmer (2013) has rekindled interest in the original line of research, added innovative quantitative evidence and further developed the concept of representational style. Grimmer notes that there are differences in the way that representatives perform their role, depending on both the characteristics of the constituency as well as the legislator herself. He focuses on how legislators connect their Washington style – the way they spend their time and resources in the capital – to their specific home style. Grimmer uses text as data methods to show that the representational style of a legislator is revealed through what she communicates to her constituents with the help of press releases. One caveat of this approach is that it assumes a monolithic constituency. Furthermore, it focuses only on the verbal. And while Fenno is very important to Grimmer, he also dismisses his version of “home style” as too focused on “folksy mannerisms”. The reason for that is because to Grimmer, home style is about information, and folksy mannerisms don’t help to inform constituents. However, to Fenno (1977), home style isn’t really about information at all. The goal that representatives, in his eyes, are pursuing, is to build trust with constituents. To this end, they face a number of challenges: On the one hand, the need to demonstrate

they are qualified to do the job. On the other hand, they need to identify and empathize with their constituents. Acts of symbolic representation are aimed at doing precisely that.

My own argument then, building on Grimmer, is that politicians have more than one representational style. They adjust their behavior according to the needs of the situation, and specifically the audience. What politicians say matters - but how they say it is also important. Representation and presentation of self through press releases is explicitly overt. Projecting likeness with constituents through subtle changes to the way spoken language is used is exactly the opposite. This strategy enables politicians to navigate the difficulties of having to satisfy a diverse set of constituents.

One Representative, Many Constituencies

I argue that politicians have diverse constituencies, meaning that their base of (potential) supporters – both voters and other people, such as lobbyists – whom they depend on, have different preferences and thus need to be treated differently. There are at least four reasons to assume this. First, constituencies are not monolithic. Many of the politicians described by Fenno (1977) have both urban and rural constituencies, and even though they generally gear themselves towards one of them, they try not to neglect the other entirely. But as shown by Cramer (2016), rural and urban constituencies can be very different, and a strategy that garners success with one leads to resentment from the other. This problem of diverse constituencies is even more pressing for senators, who represent entire states as opposed to individual districts.

Second, in order to get themselves elected, politicians have to serve more than one master. Interest groups play a substantial role in elections, whether it is through campaign finance (Desmarais, La Raja and Kowal, 2015) or the provision of information (Hall and Deardorff, 2006). A slew of research on winner-take-all politics (Hacker and Pierson, 2010) and elite domination of representation (Bartels, 2009; Gilens and Page, 2014; Achen and Bartels, 2016) has shown that economic elites exert a considerable degree of control over politicians. Acting like a ‘country boy’ may play well with the ‘folks back home’, but economic elites, who hold fundamentally different attitudes (Page, Bartels and Seawright, 2013), are likely to be less impressed by such demeanor. Politicians need to prove themselves as reliable partners to these elites, and to this end, perception is just as important as it is with voters.

Third, politicians are elites themselves. As shown by Butler (2014), representing other elites comes naturally to politicians, and doesn’t even necessitate conservative economic preferences - the presence of shared experiences is enough. Most members of Congress will likely have an easier time explaining to their constituents how to make use of school vouchers or expedite a passport, rather than how to track down a missing social security

check. This is compounded by the dynamics of Washington itself. Ultimately, “this town” revolves around green rooms, fundraisers and “\$100 haircuts” (Leibovich, 2014). If politicians employed their “folksy mannerisms” at swanky Georgetown parties, they would be laughed out of town.

Fourth, when politicians act as trustees rather than delegates (and most politicians wear both hats to at least some extent (Miller and Stokes, 1963; Saward, 2014)), voters value competence over ideological congruence (Fox and Shotts, 2009). The trustee model assumes that politicians inherently know better because they are elites - so in this case, acting like their constituents would be counterproductive for representatives.

The Shape-Shifting Representative

So politicians have some incentives to act like elites, and some incentives to act like the ‘average Joe’. Which one do they choose? According to Saward (2014), both. Relying heavily on Machiavelli, the author argues that it would be foolish for politicians to play the same role at all times. The greatest strength of his shape-shifting representative, whom he pits against the delegate and trustee models (but the argument works just as well for my purposes) is his flexibility. Rather than filling a single role, the representative of Saward (2014):

“positions him or herself as a subject with respect to constituents, supporters, or *listeners*; in other words, they adopt subject positions.” [emphasis added]

Although not as developed, this theory can already be found in Mayhew (1974), who notes that “Handling discrete audiences in person requires simple agility,” and cites a Congressman who defends his strategy of giving different speeches to pro- and anti-war crowds as following: “My positions are not inconsistent; I just approach different people differently.”

Saward (2014) follows a similar approach and goes on to note that:

“by shaping strategically (or having shaped) his persona and policy positions for certain constituencies and audiences”

the representative can have his cake and eat it too. Translated to my case, this means that depending on the audience, a politician will play either the shrewd statesman, or the ‘average Joe’. The word ‘play’ is carefully chosen here: Saward (2014) does in fact see representation as a performance - “it is performed in the theatrical sense [...] and in the speech-act sense”:

“the actor offers himself as a representative by virtue of (a) substantive policy positions, then (b) on the basis of likeness or similarity to constituents, then (c) in terms of the champion of particular interests”

If we accept that in order to get (re-)elected, politicians attempt to prove to their various constituencies that they are like them and will act in their interest – how do they do so? As pointed out by Box-Steffensmeier et al. (2003), for the concept of representation, the legislator “being like me” is at least as important to constituents as policy. After all, a roll-call vote – which is a crude indicator of legislator preferences to begin with – that is well-received in one constituency may offend another. Once a politician has voted in favor of something, they can’t take it back – every roll-call vote becomes part of their record, and can therefore be cited against them by their opponents. As noted by Maestas (2003), legislators need to take great care “to ensure that they do not cast votes or take positions that might return to haunt them.” The need to balance multiple constituencies can even go so far that legislators with more diverse constituency opinions abstain on contentious votes altogether, and thus avoid having to take a position (Cohen and Noll, 1991; Rothenberg and Sanders, 2000; Jones, 2003). The weaknesses of policy-based representation when it comes to serving the needs of multiple constituencies demonstrates the need to pay greater attention to non-policy representation.

Rhetoric and Plausible Deniability

By contrast, it is considerably easier to say one thing to one constituency, and something else to another. Grose, Malhotra and Parks Van Houweling (2015) demonstrate that senators tailor explanations of their policy positions to those of their audience, emphasizing votes and other actions that agree with a constituent’s opinion. The authors show that this strategy is effective: When there is a mismatch between a constituent’s opinion and a roll-call vote of their representative, an attenuating statement on the vote which emphasizes common ground and countervailing policy measures, improves a voter’s opinion of the senator. This strategy goes hand in hand with general ambiguity, a topic which has received greater scholarly attention. Formal-theoretic approaches show that there are considerable strategic advantages for politicians in keeping their true policy positions obscure (Shepsle, 1972; Page, 1976; Alesina and Cukierman, 1990). There is evidence that voters interpret ambiguous policy positions favorably (Tomz and Van Houweling, 2009) and that there are electoral benefits to ambiguity (Campbell, 1983).

However, even then, there is always chance that an errant comment by a politician makes it onto the record that they end up regretting later. For example, a particularly unfortunate rhetorical mishap was Romney’s infamous 47% statement, in which he

scorned (nearly) half of the voters in the country, accusing them of being unable to take care of themselves. The statement was made at a fundraiser and obviously aimed at the wealthy donors present at the event, who may have found it to be perfectly measured and reasonable. But it was an obvious slap in the face of the Republican party's considerable blue-collar base (at the time further strengthened by the Tea Party) as well as undecided voters.

There is however another, even more subtle element of communication which allows speakers to empathize with their audience. Voice¹ is much easier to modulate than any of the other means of symbolic representation described above. When speaking to a crowd of potential voters at a campaign rally or town hall event, politicians can lower their degree of articulation by dropping the g in -ing words, not releasing /t/, or speaking with a regional dialect. A Southern drawl can be code for "I am one of you" in the same way a derogatory reference to 'welfare queens' is, and it does so in a much less risky manner. When speaking to a room filled with business executives, the same politician might rely on carefully placed pauses and flawless pronunciation instead, signaling his likeness with this, very different, audience. Either way, rhetoric, generally a form of symbolic representation, thus becomes implied descriptive representation. This strategy is not just a passing fancy, but has been employed by politicians for thousands of years: Cicero (1986 English translation by J.S. Watson) the master orator of ancient Rome comments (disapprovingly) on this practice in one of his writings, *De Oratore*. He warns the budding rhetorician against the use of overly sophisticated language when in the presence of the common man, as the speaker should not wish to appear "so very wise among fools" that "though they very much approve his understanding, and admire his wisdom, yet should feel uneasy that they themselves are but idiots to him."

Not all Shape-Shifters are Alike

While Cicero's students consisted of rich, male, Roman patricians, today's crop of representatives is considerably more diverse, a fact that has been the subject of intense discussion in the representation literature. This literature centers around the impact of an increasingly diverse body of representatives, investigating questions such as the following: What does it take for women, as well as members of ethnic and sexual minorities to be elected (Sanbonmatsu, 2002; Ladam and Harden, 2017)? Once in office, how do constituencies perceive these representatives? Do voters who benefit from descriptive representation respond with a better evaluations of their representatives (Lawless, 2004; Bowen and Clark, 2014) and the institution as a whole (Scherer and Curry, 2010)? How

¹Note that throughout this paper, I use voice in the acoustic sense, as opposed to the ability of getting heard, employed by Schlozman, Verba and Brady (2012).

do these representatives vote (Cowell-Meyers and Langbein, 2009)? How do they respond differently to constituents who match them descriptively (Broockman, 2014)?

One common factor among these questions is that they ask how differences in representative identity affect their relationship with constituents *conditional* on the identity of those constituents. The question which does not get asked nearly as much is whether differences in identity lead to differences in behavior that exist irregardless of the target. Do female and minority legislators possess their own representational *style*? An example of a piece of research which does fall into this category is Niven and Zilber (2001), who study the websites with the expectation of finding a greater focus on women's issues. However, the authors find remarkably few differences between the genders, which poses the question: are there areas in which differences in style do exist?

Identity also plays an important role in the study of (socio-)linguistic variation, which has been assessed by a wide literature (see, for example, Fischer (1958); Kiesling (1998); Stuart-Smith, Timmins and Tweedie (2007); Moore and Podesva (2009); Baran (2014)). In the political arena, these issues have often cropped up in a negative manner, be it on Barack Obama receiving praise for being "articulate" (Thai and Barrett, 2007), or criticisms of the likability of female presidential candidates as a consequence of how they comport themselves with voters (Potter, 2019). A lesson to be learned from such episodes is that voters evaluate – consciously or not – whether diverse politicians conform to the rhetorical norms constituents have come to expect from straight, white, male officeholders. Whether these expectations are met or not – and whether voters respond favorably or with scepticism – the question of rhetorical and representational style as a consequence of identity merits investigating.

To sum up, I argue that politicians project likeness to demonstrate to their voters that they are like them and can be trusted. This idea revolves around symbolic acts of representation as outlined by Mayhew (1974) and even more so, Fenno (1978). I adopt Fenno's argument on the 'presentation of self' and give precedence to his focus on nonverbal communication, over Grimmer (2013)'s information-centric interpretation of home style. Furthermore, I argue that since legislators have diverse constituencies, they need to be shape-shifting (Saward, 2014), changing their appeals conditional on the audience. Policy representation is difficult in this regard, as roll-call votes commit representatives to a set position (Maestas, 2003). By contrast, rhetoric, and even more specifically, phonetic style shifting enable representatives to be equivocal about their position and adjust their representational behavior as called for by the situation.

A Theory of Phonetic Style Shifting

The theoretical claim I advance in this paper is that politicians modulate their rhetorical style in order to gain their constituents' trust. Specifically, they adjust their degree of phonetic articulation to a level demanded by, and most appropriate in a given situation. When they find themselves in need of demonstrating warmth, as would be the case on the campaign trail, they project likeness by lowering their level of articulation down to that of their constituents. When their primary goal is to establish their competence, such as in Congress, they adjust their articulation upwards, towards a more formal manner of speech. In order for a politician to succeed, they need to be both technocrat and populist, and I argue that rhetoric helps them to navigate these conflicting roles.

The role of trust in my theory is critical, and adapted directly from Fenno (1977). For Fenno, trust is important because voters cannot take the words of a politician at face value. They need to evaluate whether he a) intends to follow through on the promises he made to them, and b) has the capacity to do so. Hence, Fenno splits trust into qualification, identification and empathy. Here, I simplify Fenno's theory, as these three concepts can effectively be mapped onto the commonly used dimensions of competence and warmth employed in political psychology (Fiske et al., 2002).

So why is it that projecting warmth is more important on the campaign trail, while competence is favored in Congress? Laustsen and Bor (2017) explore the first half of this question, as they seek to adjudicate between two competing theories regarding candidate evaluations: Do voters care more about warmth or competence when it comes to making a vote choice? The authors find that – consistent with social psychology, and contrary to political science – warmth is more important. The reason here is simply that affect comes before logic – amygdala before prefrontal cortex – so that vote choice depends more on hearts than minds.

My argument for the primacy of competence in Congress rests on the notion that voters associate a certain gravity with being a member of the U.S. government (Mutz and Reeves, 2005). When in Congress, its members evidently fulfill that role, and any deviations from the social norms associated with it would be viewed negatively. Furthermore, the sociolinguistic literature has found competence-based styles to be associated with perception of greater performance in the speaker's job (Deprez-Sims and Morris, 2010). In this sense, projecting warmth would be the means of a politician getting the job, while competence is essential to doing it.

Hypotheses

My primary hypothesis can be derived directly from the theoretical discussion above: When speaking in Congress, senators attempt to demonstrate to their constituents that they are sufficiently competent to represent them. On the campaign trail, their main goal is to project warmth. Consequently, a higher level of articulation is expected in Congress.

Hypothesis 1: Articulation will be greater in congressional speeches than campaign speeches.

While the focus of this paper lies on within-legislator variation, between-legislator effects need to be considered as well. Prior scholarship in political science and sociolinguistics suggests an effect of gender. Existing research has found that female politicians are held to a higher standard by others as well as themselves (Lawless and Fox, 2004; Pearson and McGhee, 2013; Kanthak and Woon, 2015). Consequently, they would have a need to demonstrate competence more so than warmth. Furthermore, the sociolinguistics literature has consistently found that women engage in low-articulation styles of speech to a much lower degree than their male counterparts, who also shift more readily between the two styles (Fischer, 1958; Kiesling, 1998).

Hypothesis 2a: Male senators will use lower articulation in general.

Hypothesis 2b: Male senators will engage in more style-shifting than female senators, entailing a greater difference between the congressional and the campaign setting.

Data & Methods

Data Source: YouTube

The speech data used in this paper comes from the senators themselves: All U.S. senators maintain YouTube accounts on which they publish videos of their activities. Most senators use two accounts – one for their election campaigns, and one to document their legislative activities. The former contain ads, speeches at rallies, videos shot at the senators' homes (i.e. "On the Farm with Sharla and John" is an example of a video series which portrays Sen. John Tester as just an ordinary farmer from Montana) or interviews with the media. The latter focus on the legislator's legislative activities, containing sections of video from floor speeches, committee hearings, lobbying for (or celebrating) the passage of bills prioritized by the senator. This division into a campaign and a Senate account is very

convenient for me, because it divides a senator's speeches into low- and high-articulation settings, without the need for me to manually classify each video.

Consequently, I downloaded all videos from all senators from the 115th and 116th Congress who possess a campaign and a senate account.²³ This sample includes 31 Democrats⁴ and 33 Republicans. This was facilitated by the command line-based program `youtube-dl`, which enables bulk downloading of all videos on an account. To speed up the procedure, I parallelized the process through the GNU program `xargs`. Along with the videos, I also downloaded metadata and closed captions.

The next step in this pipeline consists of converting the videos from video to audio. When downloading from YouTube, videos are generally either in `.mp4`, `.mkv` or `.webm` format, with an audio sampling rate of 44100Hz (the sampling rate determines at which intervals a sample is taken from a continuous soundwave. A sampling rate of 44100Hz corresponds to 44100 data points per second). Using the Python packages `librosa` and `soundfile`, I extract the audio from each video, resample it to 16000Hz (which is more than sufficient for my purposes) and a bitrate of 256Kb/s (the bitrate determines how many possible values each data point can assume) and save it in the `.wav` format. For a more thorough explanation of these concepts, see appendix 1. Furthermore, since the videos (in particular the campaign videos) come from different sources, they vary with respect to loudness. Consequently all videos are normalized to the same base level of perceivable loudness using the program `ffmpeg-normalize`. Notably, my results vary very little between the original and the normalized audio.

At this point, I am faced with a more complicated problem. Not every second of a video consists of speech, and not every second of speech actually stems from the respective senator. The use of the senators' YouTube accounts alleviates this issue to a much greater extent than any other source, but it still exists nevertheless. To deal with this issue, I turn to two areas of machine learning and electrical/mechanical engineering: speech diarization and speaker recognition.

Speech Diarization

Speech diarization pertains to the segmentation of an audio stream into sections that actually contain speech, and the segmentation of that speech into different speakers. Importantly, speech diarization does not identify who those speakers are, it merely

²All videos uploaded until June 30, 2019, are included.

³There are also a few cases in which I use an old campaign account, because a newer one is not available. This can happen if a senator comes from an electorally secure state – they still have to build up their name recognition the first time around, but after that, they do not face the danger of being unseated and therefore forego the trouble and expense of making campaign videos.

⁴Bernie Sanders and Angus King are counted as Democrats.

separates them and keeps track of how many they are. To this end, I rely on AaltoASR, a toolkit for acoustic modeling maintained by Aalto University. This program uses Hidden Markov Models (HMM), a commonly used tool of analysis in natural language processing (both speech as well as text), to process speech features in the form of mel frequency cepstral coefficients (MFCC). See the appendix for a detailed explanation of these concepts. The output from this kind of analysis is a list of start and end times of speech, along with a form of “anonymized” identifier for the speaker, i.e. speaker 1, 2, 3, etc. wherein the repeated occurrence of the same speaker is marked as such. In addition to segmenting the speakers, speech diarization also has the advantage of filtering out unwanted noise. Applause is a very common example for this category in my dataset, and if not removed, results in a considerable disturbance of the measures later extracted from the audiostream. Diarization either cuts these sections out entirely if they are not identified as human speech, or alternatively, assigns them to their own “speaker”, which means they will be removed once speaker recognition is applied to them.

One significant downside of this method is that its computational costs increase exponentially with the duration of the video. The reason for this is that the longer the audio sample, the more different speakers and speech segments have to be tracked and squared with one another. Consequently, I only include videos that are up to 2 minutes in length. If this leads to any senator having less than 2 videos for either the congressional or campaign setting, I diarize additional, longer videos of them. One upside of limiting my sample to short videos is that it leads to a more balanced sample with respect to the campaign compared to the congressional setting. Among the 69,241 videos originally downloaded, only about 13% are campaign videos.⁵ By contrast, in the final sample, 26% are campaign videos, as they tend to be shorter, so that a higher proportion of them falls under the 2-minute mark.

Speaker Identification

The next step then consists of identifying whether any one speaker in an audio file is the respective senator or someone else. Speech by other people is quite common, for example, campaign ads, especially negative ones often have narrators and frequently contain testimonies by constituents as well. The videos of the senator in a legislative setting definitely contain a higher proportion of speech from them, but even here, there are the occasional interjections of other senators, or answers by witnesses in legislative hearings. Consequently I train a speech model for each senator using Gaussian Mixture Models (GMMs) and then compare that model with each speech sample. The training

⁵This is due to the fact that videos of congressional activities are easy to produce as they are effectively just clips from C-SPAN, whereas campaign videos are usually ads, which are expensive to make.

data was constructed as following: I hand-picked one video for each senator conforming to the following constraints: 1) Duration. The video should be around 10 minutes long. 2) Setting. The video should be from a floor speech. 3) “Purity”. The video should not contain any voice other than the senator’s in question. I limited the pool of potential videos to floor speeches because they are a very controlled setting, with few intervening noises that might confuse the model. Furthermore, each senator has spoken in this capacity, which is important for comparability. If, for example, half of my training data came from floor speeches and the other half from campaign speeches, the classifier might inadvertently learn to pick up on this distinction instead. A sufficiently long sample is also important to ensure that word choice does not influence the model. After the models are trained, I compare the log-likelihoods for each senator on every speech segment. The senator with the highest likelihood is identified as the speaker of a particular segment. All segments which are identified as stemming from the senator on whose account the video was published are marked as such and used in the subsequent analysis. Since the model is only capable of recognizing the senators, all segments of speech stemming from other people, such as constituents or committee witnesses will still be classified as senators, but generally not as the senator in question. Nevertheless, it is possible that this classification scheme involves some error, as a different person’s voice might sound closer to the senator in question than to any other senator.

Method of Analysis: Vowel Space Analysis

Given the theory of signaling likeness with constituents by shifting phonetic style, the purpose of this chapter is to test this theory in a broad sense. The goal, then, is to measure the performance of a speaker across an entire speech, or even a corpus of speeches – as opposed to specific instances of style-shifting for individual words. The linguistic concept that most closely captures the idea of sophistication and ‘elite-ness’ is articulation. I define articulation as the clarity with which phones (distinct, audible sounds that comprise speech) are produced. Here, speech can be placed on a spectrum between hyperarticulation and hypoarticulation, the former pertaining to very clear, downright exaggerated forms of articulation (e.g., like an adult might talk to a little child) whereas the latter would correspond to unclear, under-articulated speech, such as mumbling. It should be noted that my expectation is that political speech does not reach these extremes, but can be found in the continuum between them. My hypothesis, then, is that due to electoral concerns, politicians signal likeness with their constituents by modulating their overall degree of articulation in accordance with the audience. When addressing a more genteel audience, the need to signal poise and competence requires a greater degree of

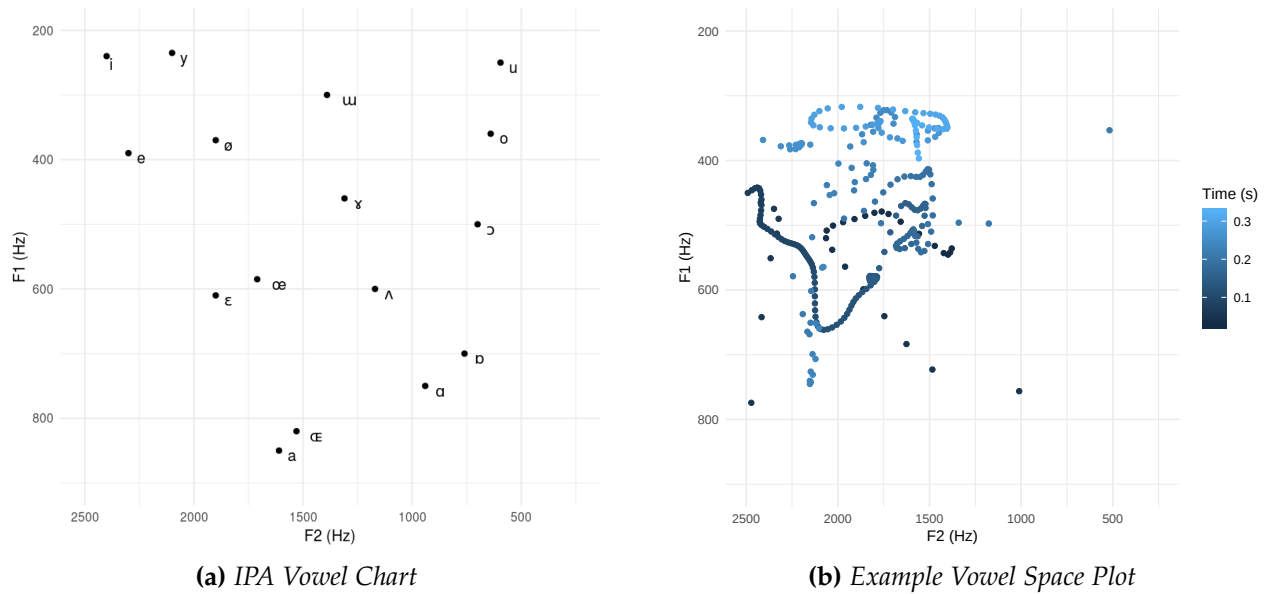


Figure 1: Figure (a) shows the International Phonetic Alphabet (IPA) vowel chart, which denominates at which F1 and F2 specific phonemes are produced on average. Figure (b) provides an example vowel space plot of the word *gettin'*, pronounced by Sen. Heidi Heitkamp. In this pronunciation, the 'i' is dropped along with the g, which is visible as there is no activity in the upper left corner where 'i' is located. The other vowel, 'e' is visible in the line of points on the left. Note that the formant range varies between every person, so values in the two plots do not overlap completely.

articulation, which means that it is closer to hyperarticulation. When addressing a more low-brow crowd, politicians need to lower their oratory sophistication, therefore moving closer to hypoarticulation. The unit of analysis here is the speech.

To measure the level of articulation in a speech, I break it down into its components – words, which themselves are comprised of phones. Between the two types of phones – consonants and vowels, the latter is far more informative of the way a speaker talks and sounds. Therefore, articulatory and acoustic phonetics largely deal with the analysis of vowels. To this end, formants are used to indicate which vocal organs are used in what way to produce a sound. The first formant (F1) corresponds to the pharynx – namely, the degree of jaw opening and the second (F2) to the oral cavity – specifically, the tongue position. For example, the vowel [i] corresponds to low F1 (tongue is high in the mouth) and high F2 (front vowel), whereas [a] yields a high F1 (tongue is low in the mouth) and low F2 (back vowel). Figure 1 (a) provides an overview of the average position of vowels - note that this can vary drastically both between and within speakers. Divergent levels of articulation will lead to different ways in which the vocal organs are used, and therefore different levels of F1 and F2. This form of analysis, the measurement of the vowel space area, is a commonly used approach in phonetics. To arrive at these measures, I follow the approach laid out by Story and Bunton (2017), measuring vowel space density.

This method largely relies on identifying how manipulation of the vocal tract leads to the production of different formants, mainly F1 (pharynx - jaw opening) and F2 (oral cavity - tongue position). The position and diversity of these formants allows researchers to make conclusions about vowel space area (VSA), with the idea that more articulate speech makes greater use of the entire VSA. In traditional VSA analysis, this is largely confined to identifying the corner vowels for very specific words, thus relying only on very small snapshots of speeches. This is useful for analyzing precisely how specific words are pronounced. Figure 1 (b) provides an example of Sen. Heidi Heitkamp pronouncing the word *gettin'*. The caveat of this approach is that it does not work as well across an entire speech. Consequently, Story and Bunton (2017) (see also Sandoval et al. (2013); Whitfield and Goberman (2014)) develop a measure they call vowel space density, plotting the use of F1 and F2 across an entire speech as a heatmap.

The first step, then, is to extract the formants. In my analysis, this is done through the program Praat and its R implementation PraatR. For this purpose, a ceiling to the formant search range of 5000Hz is set for male and 5500Hz for female speakers. First, the sound signal is down-sampled to twice that value. The audio signal is then divided into segments of 0.025s. The effective length of this window of analysis is 0.05s, because a Gaussian window is used, wherein another, tapered-off 0.0125s to each side of the central window includes signals below -120 dB. Pre-emphasis is applied, meaning that frequencies above 50Hz are amplified, wherein frequencies at 100Hz are amplified by 6dB, and another 6db for each additional 100Hz above. The purpose of this process is to enhance the signal in the noise and allow higher-frequency formants to be captured reliably. This process is illustrated in figure 6 in appendix 2. Then, the Burg LPC (Linear predictive coding) algorithm is used to calculate the frequencies for each formant. At the end of this process, a frequency value is produced for F1 as well as F2, for each 0.025s window.

The relative values of F1 and F2 pairs then provide information about which, and more importantly, how vowels were produced by the speaker. The result is converted to a two-dimensional kernel density. As described in Story and Bunton (2017), the density values are normalized to a range of [0,1] by dividing each value by the maximum value. This process ensures greater comparability between speeches and speakers. Finally, a convex hull is drawn around the area with a normalized density of 0.25 or higher, indicating the speaker's vowel space area across the speech. The area enclosed within this hull, calculated with the package *phonR*, is my indicator for the level of articulation of a speech, and can then be compared to other speeches. Since the area is large and varies greatly, I use the log of this value in the subsequent analysis. Figure 2 compares two such densities for speeches of Senator Heidi Heitkamp – one from a campaign ad, and the other from

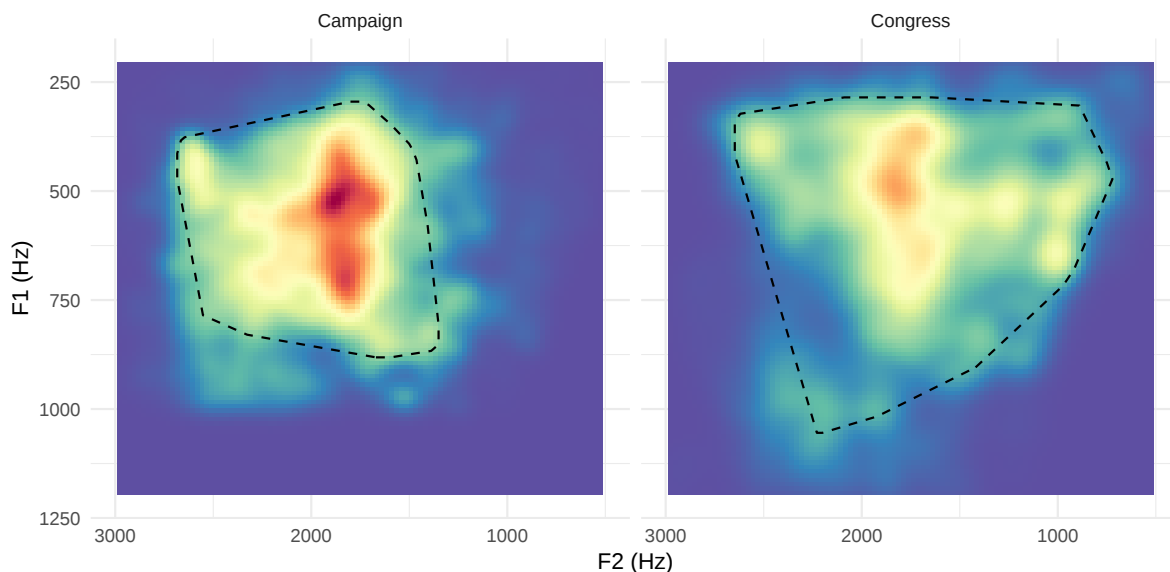


Figure 2: *Vowel space plot for two concatenated speeches of Sen. Heidi Heitkamp – one in a campaign context, and one in Congress. A larger area covered by the high-density areas (enclosed in the black dashed line) corresponds to a greater degree of articulation.*

Congress. The black dashed line illustrates the convex hull, which encompasses a greater area in the congressional speech.

Dataset

69,241 videos were originally downloaded. Of these, 12,509 were short enough to be diarized. Since many of these do not contain speech by a senator (for example because they are narrated campaign ads, or because the speaker recognition classifier isn't sufficiently confident), this number is further reduced to 4426. This number decreases to 4078 after the formant detection, as the algorithm fails to detect F1 and F2 in some cases. At this point, I also exclude outliers with respect to log vowel space area, removing all videos that fall under the 5th and over the 99th quantile (I remove more on the lower end as the distribution is left-skewed). Furthermore, I remove all senators who do not have at least two videos of each type left after this process. In the end, I am left with 3466 videos,⁶ running for a combined 54 hours. This sample includes 64 senators – 31 Democrats and 33 Republicans. Of these, 13 are women, 23 are from the South, and 21 come from states where at least 30% of the population lives in rural areas.

⁶3429 videos for the non-normalized audio. The numbers differ because the formant detection works slightly better on the normalized audio.

Audio Data in Political Science

Having detailed my own methodology, I now provide an explanation of how it relates to other approaches employed in audio-based research in political science. A number of research projects using audio data have emerged in recent years: Hwang, Imai and Tarr (2019) attempt to emulate the human work of the Wisconsin Advertising Project/Wesleyan Media Project through machine learning tools, with the goal of transcribing the audio, summarizing the issues being discussed, coding the level of negativity and detecting whether the opponent is mentioned (as well as other, purely image-based concepts such as face detection and recognition, as well as image text recognition). Rheault and Borwein (2019) rely on both image and audio features to code emotion in political videos. Knox and Lucas (2019) treat speech as a physical realization of latent political concepts, such as emotion, emphasis and issue positions. This permits them to construct a general-purpose model of speech, through which these concepts can be approximated. Dietrich and Mondak (2019) is somewhat similar in the sense that they use audio features derived from phone interviews to predict vote choice. Dietrich, Enos and Sen (2018) and rely on pitch to measure emotional arousal in oral arguments in the U.S. Supreme Court, which they again use to predict vote choice. Similarly, Dietrich, O'Brien and Jielu (2019) also use pitch to measure emotion, and combine this data with word-level timing data, derived from the use of forced alignment on speech transcripts.

A commonality of these papers – and mine – is that they all rely on features derived from a spectral analysis of the frequency band. The differences lie in how these features are used. Here, I explain these commonalities and differences in a technical sense. In the discussion section (see below), I deliberate on the philosophical differences underlying the technical ones. Now, a number of different features can be extracted from a spectral analysis of audio data. Among these are pitch, MFCC features (see Appendix 1), and formants, none of which are used exclusively by only one set of authors in this literature. In fact, pitch and formants are effectively two sides of the same coin – another term for pitch is F0, as it is the fundamental frequency, and calculated in a manner very similar to F1 and F2 (and to do so, Dietrich, Enos and Sen (2018) and I also use the same program). One distinguishing point between these papers is how many different types of features are used. On the surface, this seems like a fairly mundane, technical matter, but it is indicative of the purpose for which they are used. Hwang, Imai and Tarr (2019), Rheault and Borwein (2019) and Knox and Lucas (2019) all use a very large number of different features (in fact, one of the characteristics that distinguishes Knox and Lucas (2019) from even the computer science literature is their number of features), and many of these features are very high-dimensional, such as MFCCs. The reason for this is that the goal of

these studies is *prediction*, be it of negativity, emotion, or content. Here, a larger number of features leads to higher accuracy. The exact relationship between each feature and the outcome variable is of no interest – there is no expectation that, say, the 19th MFCC component would be particularly good at predicting, for example, anger. This is the crucial difference to my own study (and some of the work of Dietrich and colleagues, in particular Dietrich, O’Brien and Jielu (2019)).⁷ Formants are not just a set of abstract audio features that happen to be good at predicting something – they directly relate to the process of human speech production, namely the position of the tongue. In turn, the position of the tongue is directly connected to vowels, the main building block of human language. As such, my approach is more deductive, and is connected more closely to the linguistics and phonetics literatures, which have researched formants for decades. In addition to the quantity of interest in my analysis, my application of speaker diarization – the segmentation of an audio stream into different speakers – is so far also unique in political science.

Model Specification

With a measure for vowel space of each speech in hand, the main hypothesis can be put to the test. To reiterate, I expect greater informality in speeches given in a campaign context, compared to a legislative setting. Consequently, I use a linear model where the unit of analysis is a speech. The **dependent variable** is **log vowel space area**. A dummy indicating whether a speech was given in a **congressional or a campaign setting** constitutes the **main independent variable**. An important property of my data is that the individual speeches are clustered by senators, each of whom have their own individual style of speaking. Vowel space area cannot be directly compared between individuals – comparisons must take place *within* senators. Generally speaking, this would entail a within-unit, fixed effects model. However, my model specification also includes a number of variables (gender, party, South – see below) which do not vary within senators and are therefore fully collinear with the fixed effects. Consequently, such a model would not be identifiable. Instead, I rely on a **mixed-effects model with a random intercept** for each senator. This approach retains as much of the desirable statistical properties of a pure fixed-effects model as possible, while still being fully identifiable.

While my primary interest lies in the effect of the setting on articulation, there are a number of other variables to consider. First among them is *gender*. There are two reasons

⁷However, it should be noted that in the part of my study where prediction is important – speech diarization and speaker identification, a large number of features, among them MFCCs, are used. The same is true for the forced aligner employed by Dietrich, O’Brien and Jielu (2019).

for this: One, gender is of theoretical interest, as discussed above. Much recent scholarly work has gone into revealing how female legislators differ from their male colleagues, both in their behavior and in how voters respond to them. To reiterate, under hypothesis 2a, I expect male senators to have a lower degree of articulation. Hypothesis 2b predicts greater style-shifting for male senators. To account for the latter, I introduce an interaction effect between *gender* and the *campaign setting* dummy. Second, if nothing else, gender is needed as a control since it plays a large role in speech production.

Additional variables are in the model primarily as controls, but might also be of theoretical interest to scholars. The Republican party is geared towards appealing to rural voters (and the sociolinguistic literature identifies the urban/rural cleavage as one of the largest sources of speech variation) and also attempts to project a strong “patriotic”, pro-American image. Consequently, it is plausible that Republican senators project these values through lower articulation. On the other hand, the Republican party also appeals to the wealthy, which would lead to the opposite expectation. Either way, *party* is sufficiently relevant to be controlled for. For the same reasons, I also control for the proportion of the *rural population* within the state directly. With a Pearson’s R of 0.09,⁸ the correlation between the party and the rural variable is not as large as might be expected, so there is clear merit in having both variables in the model. Furthermore, I include a dummy indicating whether a senator is from the *South*. The special role of the South in American Politics has been documented extensively (Key, 1949) and since the region also possesses its own speech patterns, it is to be expected that Southern senators use a lower level of articulation. Finally, speech *Duration* (in minutes) is used purely as a control. While the use of a density for vowel space is already intended to account for the length of the speech, introducing this variable into the model ensures that this potential confounder is properly and fully controlled for.

Results

The full results for this model can be found in Appendix 2, Table 1. Henceforth, I will be referring to the results from model 1 in this table, which contains the full specification, based on the most reliable version of the audio data.⁹

⁸Since my sample includes more videos of some senators, this correlation is partially the result of a somewhat unbalanced sample. In a sample with only one data point per senator included in this study, the correlation is 0.21

⁹Models 1 and 2 describe the results for audio data for which the level of loudness was normalized, whereas models 3 and 4 cover the original audio. Differences between the two approaches are minimal. While the interaction term between setting and gender is important for both theoretical and practical reasons, it does make the interpretation of the results with respect to setting a little less straightforward. Therefore, models 2 and 4 show a version of the results without this term. While there are some differences

The results are as following: As expected, the relationship between vowel space area and the campaign setting is negative and significant. This finding – that speaking on the campaign trail is associated with lower articulation than in Congress – provides evidence in favor of hypothesis 1. Specifically, a change in the setting variable from Congress to corresponds to a marginal effect of -0.17 for female and -0.02 for male senators. Given that the mean of the log vowel space area is 12.92, with a standard deviation of 0.49, this entails a 0.35 standard deviation vowel space area reduction for female senators and a 0.04 standard deviation vowel space area reduction for their male counterparts. In addition to supporting hypothesis 1, these numbers also provide evidence *against* hypothesis 2b, as the effect for changing setting is larger for women (a 0.35 standard deviation change, compared to 0.04) than for men. This result can also be found in Figure 3.

The relationship between gender (male) and vowel space area is also negative and significant. This is expected, as specified in hypothesis 2a. Substantively, male senators speak with a lower degree of articulation than female senators. To be precise, the marginal effect of -0.12, corresponding to going from female to male, in a campaign setting, corresponds to a 0.24 standard deviation reduction of vowel space area. In the congressional setting, these numbers are -0.28 and 0.57, respectively. This also means that the difference between male and female senators is larger in Congress than on the campaign trail. Figure 4 also shows this result.¹⁰

The party variable – while not my primary interest in this study – also yields an interesting result (see Figure 5). A change in party from Democrat to Republican corresponds to a reduction in log vowel space area of 0.15. This means that Republicans speak with a lower degree of articulation than Democrats. At the same time, the estimates for proportion of rural population in the state, as well as South, are both statistically insignificant. A likely theoretical interpretation of this result is that the Republican party has become synonymous with rural and Southern identity, and therefore absorbs the reduction in vowel space that might otherwise be expected from these variables. Finally, the estimate of the video duration variable is not statistically significant, and the effect size is very small. This suggests that the use of a density to account for variance in speech length already does a good job in controlling for this potential confounder on its own.

in the effect sizes, these two models nevertheless confirm to the overall results of the full models.

¹⁰Note that while the confidence intervals overlap, Table 1 in Appendix 2 shows that this relationship is nevertheless statistically significant.

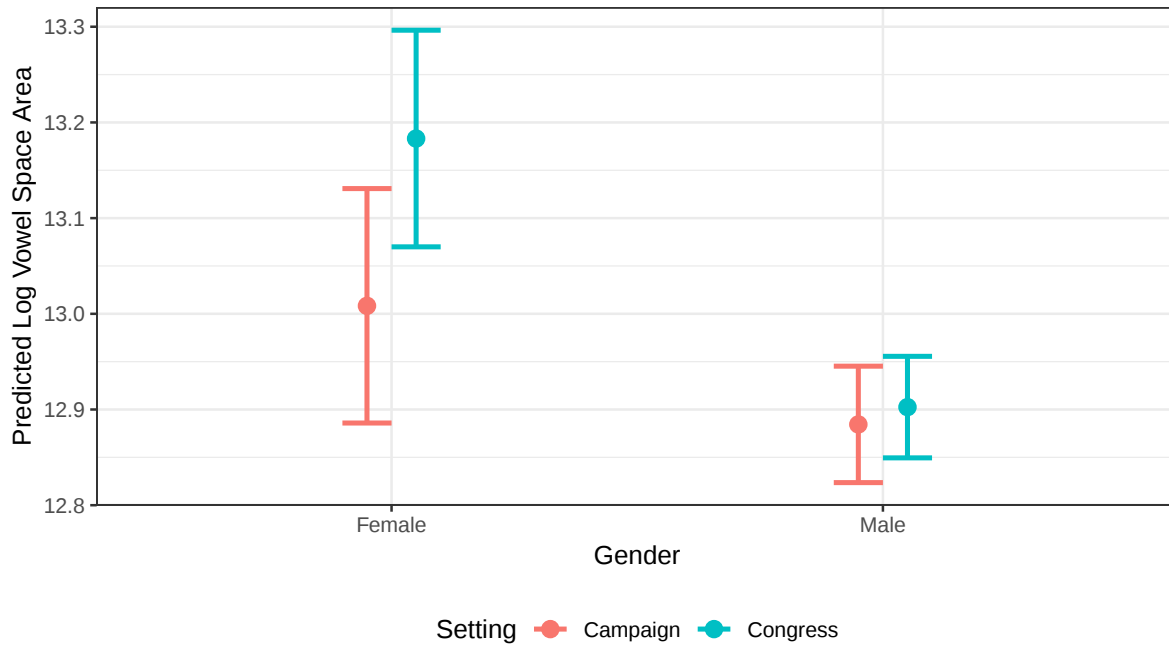


Figure 3: Predicted log vowel space area in a given setting, conditional on gender. The plot shows that legislators use a greater degree of articulation in Congress than on the campaign trail, and that this difference is larger for female legislators.

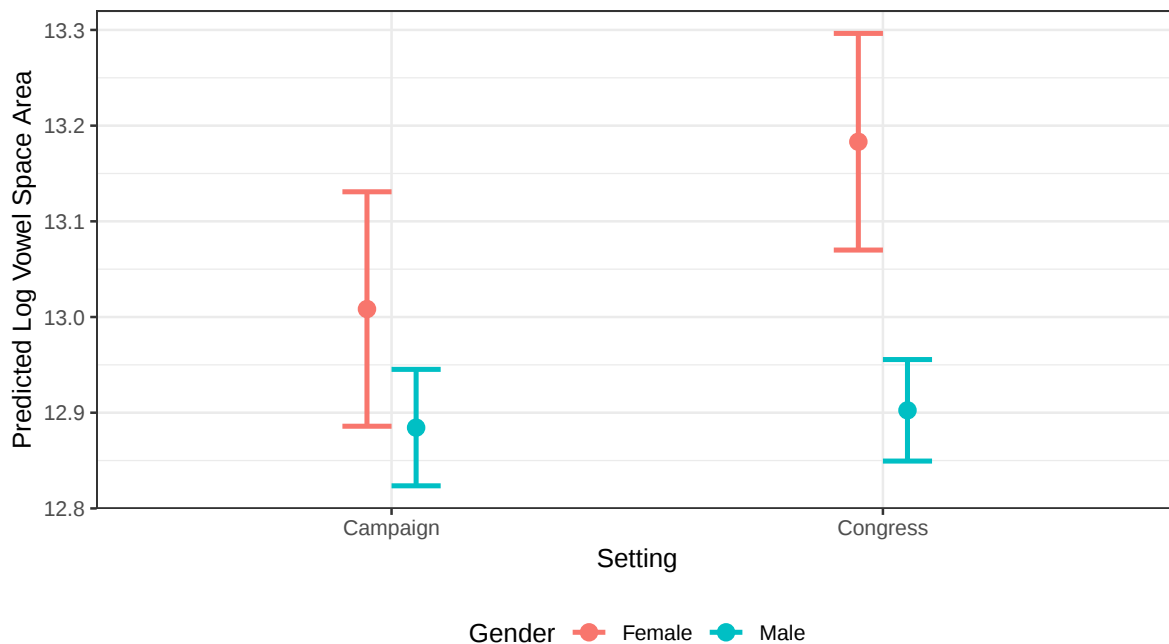


Figure 4: Predicted log vowel space area for a given gender, conditional on the setting. The plot shows that female legislators use a greater degree of articulation than their male counterparts in general, and that the difference between the genders is larger in congress than on the campaign trail.

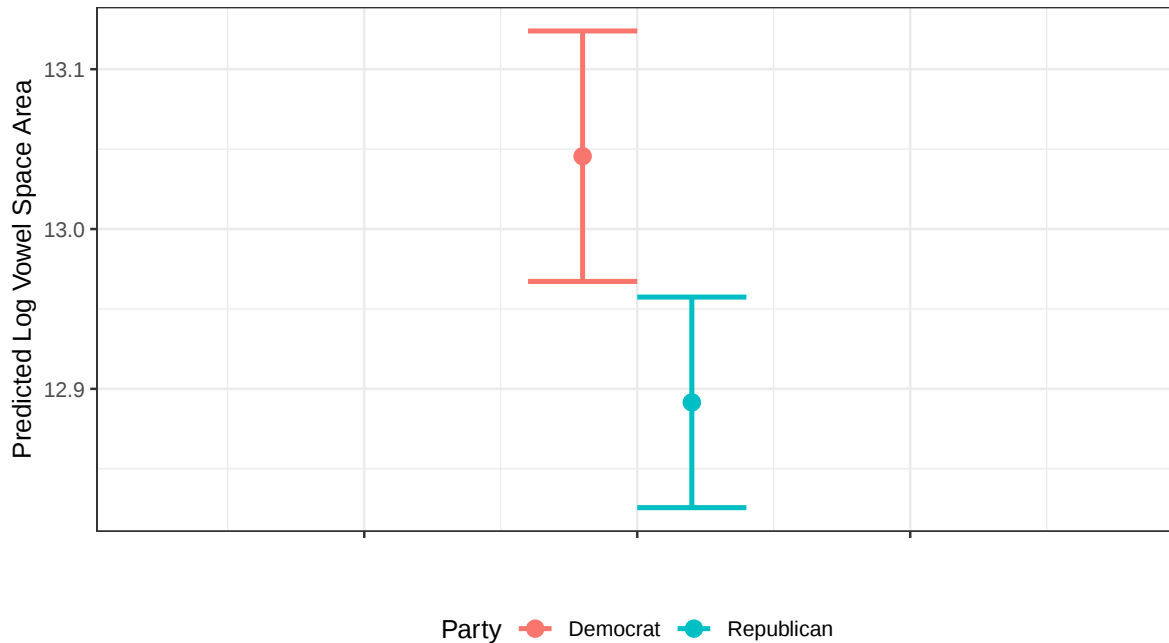


Figure 5: *Predicted log vowel space area for a given party. The plot shows that Democratic legislators use a greater degree of articulation than their Republican counterparts.*

Discussion

In this paper, I have presented evidence for the theory that politicians engage in phonetic style-shifting in order to project likeness with an audience. Electoral goals drive this phenomenon, as representatives find themselves in the position of having to appeal to multiple constituencies, with potentially conflicting interests and opinions. The analysis presented here shows that politicians do indeed speak in a more high-brow manner when fulfilling their legislative role, as indicated by a greater vowel space area. It seems plausible that this behavior is driven by the need to demonstrate competence and impress the political sophisticates. By contrast, office-holders can demonstrate warmth by addressing voters in a more colloquial tone.

This phenomenon tells us something about representation. My findings match the conclusions of Fenno (1977) more closely than those of Grimmer (2013). Politicians do not play either the ‘statesman’ or the ‘appropriator’, at the detriment of the other. Rather, as predicted by Fenno, politicians adopt a specific style in accordance with their audience, and they are capable of switching between these styles fluidly, as posited by Saward (2014). Grimmer (2013) does touch on the notion that some senators with multiple constituencies, such as Hillary Clinton, focus on two areas (in this case pork and policy). But as discussed above, this carries significant risks because it often means going on the

record with positions that, while pleasing one constituency, will offend another. It is perhaps no coincidence that the most recent losers in presidential elections (Clinton and Romney most prominently, but to a lesser extent also McCain and Kerry) were known for attempting this balancing act, but failing to do so convincingly, thus coming off as insincere “flip-floppers”. The approach analyzed in this paper – phonetic style-shifting – affords the practitioner much greater plausible deniability, and, thanks to how finely calibrated humans are to communicative subtleties, might be just as effective.

My findings with respect to gender merit additional discussion. On the one hand, I find that female legislators make greater use of their vowel space, which entails a greater level of formality. This is consistent with my expectations. Prior research has shown that female representatives are more qualified than their male counterparts (Pearson and McGhee, 2013), but also held to a higher standard (Lawless and Fox, 2004; Kanthak and Woon, 2015). Consequently, they put more emphasis on their qualifications than male legislators (Niven and Zilber, 2001). Greater formality in speech across the board is consistent with this style, as it emphasizes competence. However, the finding that women representatives style-shift more is somewhat unexpected. Figure 3 shows that the difference between the campaign and congressional setting is much greater for women than for men. This finding stands in contrast to the sociolinguistics literature, which expects males to style-shift more. Apparently there is something about politics that reverses this relationship. The most plausible explanation for this phenomenon is the ‘likeability’ conundrum female candidates seem to be facing. On the one hand, there is a need for them to appear tough, smart and competent, hence the overall higher level of articulation. But on the other hand, voters want their representatives to be authentic and relatable (Fenno, 2007), and they judge women more harshly if they fail to do so. It appears that female representatives respond to this problem by lowering their degree of articulation by a much greater amount when going from Congress to the campaign trail. At this point, it should be noted that given the implicit connections of race and language, a similar effect might exist for legislators belonging to ethnic minorities. I have not tested this relationship in this paper, as the combined number of Hispanic, African-American and Asian-American senators in the 116th Congress is only 9 (with Kamala Harris in the latter two categories) – too low for a rigorous statistical analysis. This gap could be addressed in future research, for example by relying on larger samples from the House or state legislatures.

The findings in this paper also raise questions about the role of policy in representation. The traditional view of representation largely defines it as a principal-agent relationship, where legislators implement the policy preferences of their constituents (APSA, 1950; Schattschneider, 1960). The recent scholarly literature however has shown that reality

does not match this theory. For one, it appears that it is the citizens who adapt policy preferences from their representatives, and not the other way around (Lenz, 2009; Broockman and Butler, 2017). Consistent with this finding, representatives do not seem to care what the preferences of their constituents are (Kalla and Porter, 2019). Furthermore, voters appear unwilling to remove an incumbent representative with diverging policy preferences unless the mismatch is particularly egregious (Hollibaugh, Rothenberg and Rulison, 2013). These findings undermine the foundations of the constituent-representative relationship. Given this background, the concept of likeness, presented in this paper, provides representatives with an alternative means of building a relationship with their constituents for the purpose of getting re-elected. Likeness through style-shifting is intended to convince voters that their representative is on their side. If we believe that this strategy is effective with voters – and Ricks (2018) provides evidence that it is – then this has implications for representative behavior. Rather than becoming policy wonks adept at passing legislation in line with the preferences of their constituents, representatives would be incentivized to invest time and resources into more appearance-focused activities. Given trends such as the increasing proportion of congressional staffers assigned to communication (Lee, 2016), the growing importance of appearances on the talk show circuit (Baum, 2005; Leibovich, 2014), and the use of Congressional hearings for making TV-ready soundbites rather than fact-finding (Greenfield, 2019), empirical reality does seem to match this expectation. In light of this evidence, it is worth asking whether for the study of representation, the relative importance of political communication, compared to policy outcomes, should be re-evaluated.

The findings of this paper and the theory of likeness also speak to the study of populism. Populism is part of the political fabric of many countries, not least of all the United States. And while populism comes in many guises, one unifying feature of its leaders is that they claim to speak for “the people” (Kazin, 1998). Likeness as a representational style is very similar in this sense, as it also aims to address “the people” and doesn’t require any policy compromises for their representation. Furthermore, populism contains a gender component – it often relies on appeals from male politicians to male voters (Spierings and Zaslove, 2017; Bobba et al., 2018). Here, I find that male senators use a lower level of articulation – and thus, essentially, a greater appeal to the ‘common man’. Of course, likeness and populism are not the same in every way. While populists generally (claim to) fight against the establishment, a likeness-based presentation of self is not limited in this way. If anything, I claim the opposite – politicians who are often firmly part of the establishment, rely on style-shifting to maintain the support of voters whom they do not represent in terms of policy. Consequently it would be more appropriate to consider likeness-based presentation of self a strategy, which can be

used by populists and members of the establishment alike. In this sense, likeness-based representation is more closely related to personalism (Kostadinova and Levitt, 2014) than populism. Furthermore, my research also speaks to the study of populism due to its focus on the analysis of rhetoric – a very prominent feature of this literature (Moffitt, 2016). However, this branch of research (at least in terms of quantitative work) has focused almost exclusively on the verbal component of rhetoric. The content of populist speech may be fascinating to scholars and journalists due to its shock value – but it isn’t necessarily what makes it effective. Populism relies in no small part on charismatic leaders – and charisma is more about delivery than content. For example, Klostad, Anderson and Peters (2012) show that political leaders with lower-pitched voices tend to be rated as more appealing. Future research in this vein could greatly contribute to our further understanding of populism.

The research carried out in this paper represents one of the first attempts in political science to leverage the information contained in the audio of speeches. Other studies, such as Hwang, Imai and Tarr (2019), Rheault and Borwein (2019), Knox and Lucas (2019), Dietrich, Enos and Sen (2018), Dietrich, O’Brien and Jielu (2019), and Dietrich and Mondak (2019) follow the same basic principles of extracting information from audio data, but use different sets of features, and more importantly, use these features for different purposes. On the one hand, there is the use of audio for the purpose of prediction. This can take shape in the form of tasks that are still closely connected to audio and human speech – such as speaker diarization, identification and speech transcription – or the prediction of more latent concepts, such as emotion, which can also be measured with (or in combination with) other types of data, such as images and text. But in addition to being useful for prediction, audio data also conveys meaning on its own. My own, deductive approach, relying on formants to measure a speaker’s vowel space area, is more directly connected to the process of human speech production in itself, as well as the fields – (socio-)linguistics and phonetics – that possess the most expertise on it. The connection with sociolinguistics specifically also yields a particular advantage due to its status as a social science: theory. Sociolinguistics have spent decades studying why human speech varies in terms of pronunciation and articulation, be it with respect to the situation, the conversational partner, the speaker themselves, and so on. This theory can be employed to generate predictions which can be tested for their applicability in the political arena. Machine learning, by contrast, has no such priors to offer.

But of course, even my formant-based approach still carries some predictive component, as humans cannot identify formants directly, just as much as they are unable to detect MFCCs. So perhaps, it makes more sense to view this distinction as a continuum, on which the different approaches can be arranged. On the pure machine learning-based end

would be the applications of Hwang, Imai and Tarr (2019) and Rheault and Borwein (2019), with Knox and Lucas (2019) a little further along the line. My own research sits at the other end of this spectrum, as it is much more closely connected to linguistics and phonetics. Pitch-based approaches (mostly by Dietrich and colleagues) lie somewhere in the middle of this continuum, as pitch has useful predictive properties, but also carries substantive meaning on its own. My goal here is not to argue that one of these approaches is better. Instead, it depends on the application. If the goal is to study political communication for its own sake, the insights the linguistics community has gained over the decades should not be ignored. By contrast, if the purpose is to study how communication reflects on some other, latent, concept, the predictive approach has its advantages. Of course research design and data availability (including training data) also play a crucial role. Furthermore, while my analysis of vowel space area is based on linguistics and phonetics, my application of speaker diarization and identification has its roots in computer science and machine learning – so there is no reason to expect the two cannot be combined. That being said, for future research, I expect more development in political science on the linguistics end of this spectrum, as the connection between these two social sciences provides us with a competitive advantage we do not enjoy in machine learning.

A correlate of the novelty of audio-based research designs in political science is that they necessarily cannot address every potential issue and therefore come with a some of caveats. This is, of course, also true for this research: For one, the dichotomy of campaign versus legislative YouTube accounts is helpful in practice, but it may be overly simplistic. One, not every senator keeps to this practice. A few senators use only one account, and post both campaign and congressional activities on it. Manual coding of these videos into either category would allow for a moderate increase in sample size.

By not hand-coding anything about the setting of the video, I am also unable to determine which exact context it took place in. For example, I do not have any information on whether a video steams from a campaign ad, a rally or a town hall. Similarly, it might be of interest whether a campaign video was an ad which might have been broadcast in an area with a specific socio-demographic profile.

On a theoretical level, there is another question: do senators engage in style-shifting intentionally, or is this an automatic and involuntary response? There is little doubt about the notion that on some level, all humans engage in situational style-shifting involuntarily and inevitably (Schilling-Estes, 2003; Campbell-Kibler, 2009, 2010). However, this does not mean that politicians can't train themselves to become better at it. Nunberg (2008) reports on an interesting example: When Sarah Palin originally entered the political stage, she engaged in g-dropping by default. As she became more experienced, she learned to eliminate this folksy speech variant in order to sound more professional. Her further

evolution as a politician led her to re-introduce g-dropping later in her career, but only when the situation called for it. This learning curve – from involuntary informal speech, to more formal language in general, and finally, in the ability to control the style-shifting as needed, does illustrate that political actors can learn to control this strategy. This also speaks to an evolution of political professionalism. While the scholarly literature largely focuses on professionalization in the legislative arena (Berry, Berkman and Schneiderman, 2000; Carsey, Winburn and Berry, 2017), the proliferation of advisors, PR gurus and public speaking trainers has also allowed political professionals to polish the manner in which they present themselves to the public.

From a representative point of view, it is also fairly inconsequential whether a senator reminds herself before each speech to speak in the appropriate style, or does so involuntarily. My purpose here is not to measure how “insidious” politicians are in their “deception” of voters – but to analyze a strategy. Even if it is entirely automatic, it is still worth knowing how it happens, who is good at it, and to explore how it matters for representation. For its effectiveness, and its implications for symbolic and descriptive representation, it is not important whether it is intentional, or simply comes natural to (some) politicians. The degree to which style-shifting is involuntary is an interesting question worth researching. But before we can get to this point, we first need to establish the extent to which it exists – and this paper does that. The next step then, is to ponder the why.

The theory and empirical results presented in this paper raise further questions which I plan to address in future research. While vowel space area is a frequently used measure from the phonetics literature which neatly captures the concept I am trying to measure, it is not the most intuitive approach. To non-phoneticians, a vowel space area of, say, 12, does not really mean anything. Furthermore, I have largely motivated my research with common examples of “folksy mannerisms” such as g-dropping. Consequently I plan to develop a method for the detection of this kind of phrasing and test whether it relates to the target audience in a similar manner. Research done by Yuan and Liberman (2011) has already approached the issue of measuring g-dropping directly, and I intend to build on this literature.

Another aspect of this research, which is directly related to my theory, is the question of topic. For example, issues such as unemployment, crime or immigration have a very direct appeal to ordinary citizens. By contrast, other issues, such as foreign policy are further removed from the common man and therefore carry more interest for political sophisticates (this also touches on the debate originating from, among others, (Miller and Stokes, 1963). It follows that depending on the issue they are currently discussing, politicians would assume either a low- or high-brow style of talking. Consequently I

intend to measure the within-speech variance in articulation, conditional on the topic. This also connects my research to the other area of natural language processing which is far more prominent in political science - text.

Finally, all of this research *assumes* that politicians modulate the degree of articulation in their speech because different forms of speaking have different effects on the audience. However, causal evidence for this relationship is another matter. Through a survey experiment in Thailand, Ricks (2018) has shown that different styles of speaking do indeed have the expected effect on listeners: While informal and local language cause respondents to rate politicians higher on likability and kinship, formal language is seen as a signal for competence. I plan to carry out a similar experiment in the U.S.

This paper then represents the first exploration of this topic - it establishes a theory of phonetic style-shifting in the service of symbolic and implied descriptive representation. Furthermore, it tests that theory by developing a measure for the primary theoretical concept in question - articulation. This measure – vowel space area – permits me to test the theory in the broadest sense possible, but it also sets me up for investigating its subtler aspects. The evidence presented here supports the theory of rhetorical style shifting, and I plan to expand on it in future work.

References

- Achen, Christopher and Larry M. Bartels. 2016. *Democracy for Realists: Why Elections Do Not Produce Responsive Government*. Princeton: Princeton University Press.
- Alesina, Alberto and Alex Cukierman. 1990. "The Politics of Ambiguity." *The Quarterly Journal of Economics*, CV(4):1998–2006.
- APSA. 1950. "Part I. The Need for Greater Party Responsibility." *The American Political Science Review* 44(3):15–36.
- Arnold, R. Douglas. 1990. *The Logic of Congressional Action*. Yale University Press.
- Baran, Dominika. 2014. "Linguistic Practice, Social Identity, and Ideology: Mandarin Variation in a Taipei County High School." *Journal of Sociolinguistics* 18(1):32–59.
- Bartels, Larry M. 2009. *Unequal Democracy: The Political Economy of the New Gilded Age*. Princeton: Princeton University Press.
- Baum, Mathew A. 2005. "Talking the Vote: Why Presidential Candidates Hit the Talk Show Circuit." *American Journal of Political Science* 49(2):213–234.

- Berry, William D., Michael B. Berkman and Stuart Schneiderman. 2000. "Legislative Professionalism and Incumbent Reelection: The Development of Institutional Boundaries." *The American Political Science Review* 94(4):859–874.
- Bobba, Giuliano, Cristina Cremonesi, Moreno Mancosu and Antonella Seddone. 2018. "Populism and the Gender Gap: Comparing Digital Engagement with Populist and Non-populist Facebook Pages in France, Italy, and Spain." *International Journal of Press/Politics* 23(4):458–475.
- Bowen, Daniel C. and Christopher J. Clark. 2014. "Revisiting Descriptive Representation in Congress: Assessing the Effect of Race on the Constituent–Legislator Relationship." *Political Research Quarterly* 67(3):695–707.
- Box-Steffensmeier, Janet M., David C. Kimball, Scott R. Meinke and Katherine Tate. 2003. "The Effects of Political Representation on the Electoral Advantages of House Incumbents." *Political Research Quarterly* 56(3):259–270.
- Broockman, David E. 2014. "Distorted Communication, Unequal Representation: Constituents Communicate Less to Representatives Not of Their Race." *American Journal of Political Science* 58(2):307–321.
- Broockman, David E. and Daniel M. Butler. 2017. "The Causal Effects of Elite Position-Taking on Voter Attitudes: Field Experiments with Elite Communication." *American Journal of Political Science* 61(1):208–221.
- Butler, Daniel M. 2014. *Representing the advantaged: How politicians reinforce inequality*. Cambridge University Press.
- Campbell, James E. 1983. "The Electoral Consequences of Issue Ambiguity: An Examination of the Presidential Candidates' Issue Positions from 1968 to 1980." *Political Behavior* 5(3):277–291.
- Campbell-Kibler, Kathryn. 2009. "The nature of sociolinguistic perception." *Language Variation and Change* 21(1):135–156.
- Campbell-Kibler, Kathryn. 2010. "The sociolinguistic variant as a carrier of social meaning." *Language Variation and Change* 22(3):423–441.
- Carsey, Thomas M., Jonathan Winburn and William D. Berry. 2017. "Rethinking the Normal Vote, the Personal Vote, and the Impact of Legislative Professionalism in U.S. State Legislative Elections." *State Politics and Policy Quarterly* 17(4):465–488.

- Cicero, Marcus Tullius. 1986. *De Oratore*. J.S. Watson (Translator).
- Cohen, Linda R. and Roger G. Noll. 1991. "How to vote, whether to vote: Strategies for voting and abstaining on congressional roll calls." *Political Behavior* 13(2):97–127.
- Converse, Philip E. 1964. The Nature of Belief Systems in Mass Publics. In *Ideology and Discontent*, ed. David Apter. University of Michigan.
- Cowell-Meyers, Kimberly and Laura Langbein. 2009. "Linking Women's Descriptive and Substantive Representation in the United States." *Politics & Gender* 5(2009):491–518.
- Cramer, Katherine J. 2016. *The politics of resentment: Rural consciousness in Wisconsin and the rise of Scott Walker*. University of Chicago Press.
- Deprez-Sims, Anne Sophie and Scott B. Morris. 2010. "Accents in the workplace: Their effects during a job interview." *International Journal of Psychology* 45(6):417–426.
- Desmarais, Bruce, Raymond J. La Raja and Michael S. Kowal. 2015. "The fates of challengers in U.S. house elections: The role of extended party networks in supporting candidates and shaping electoral outcomes." *American Journal of Political Science* 59(1):194–211.
- Dietrich, Bryce J, Diana Z O'Brien and Jielu. 2019. "Do Representatives Emphasize Some Groups More Than Others?".
- Dietrich, Bryce J and Jeffery J Mondak. 2019. "Using the Audio from Telephone Surveys for Political Science Research Early Draft and Preliminary Results : Please do not circulate without permission.".
- Dietrich, Bryce J., Ryan D. Enos and Maya Sen. 2018. "Emotional Arousal Predicts Voting on the U.S. Supreme Court." *Political Analysis* pp. 1–7.
- Downs, Anthony. 1957. "An Economic Theory of Political Action in a Democracy.".
- Eulau, Heinz and Paul D. Karps. 1977. "The Puzzle of Representation: Specifying Components of Responsiveness." *Legislative Studies Quarterly* 2(3):233–254.
- Fenno, Richard. 1978. *Home style: House members in their districts*. New York: Pearson College Division.
- Fenno, Richard. 2007. *Congressional travels: places, connections, and authenticity*. Routledge.
- Fenno, Richard F. 1977. "U.S. House Members in Their Constituencies: An Exploration." *American Political Science Review* 71(3):883–917.

- Fischer, John L. 1958. "Social Influences on the Choice of a Linguistic Variant." *WORD* 14(1):47–56.
- Fiske, Susan T., Amy J.C. Cuddy, Peter Glick and Jun Xu. 2002. "A model of (often mixed) stereotype content: Competence and warmth respectively follow from perceived status and competition." *Journal of Personality and Social Psychology* 82(6):878–902.
- Fox, Justin and Kenneth W. Shotts. 2009. "Delegates or Trustees? A Theory of Political Accountability." *The Journal of Politics* 71(4):1225.
- Gilens, Martin and Benjamin I. Page. 2014. "Testing Theories of American Politics: Elites , Interest Groups , and Average Citizens Martin Gilens Benjamin I . Page forthcoming Fall 2014 in Perspectives on Politics." *Perspectives in Politics* 12(3):564–581.
- Goffman, Erving. 1959. *The Presentation of Self in Everyday Life*. New York: Doubleday.
- Greenfield, Jeff. 2019. "The Michael Cohen Hearing Didn't Have to Be So Awful."
URL: <https://www.politico.com/magazine/story/2019/02/27/michael-cohen-hearing-225343>
- Grimmer, Justin. 2013. *Representational Style in Congress: What Legislators Say and Why It Matters*. Cambridge University Press.
- Grimmer, Justin Ryan. 2010. "Representational Style: The Central Role of Communication in Representation." *PhD Thesis* .
- Grose, Christian R., Neil Malhotra and Robert Parks Van Houweling. 2015. "Explaining Explanations: How Legislators Explain their Policy Positions and How Citizens React." *American Journal of Political Science* 59(3):724–743.
- Hacker, Jacob and Paul Pierson. 2010. *Winner-take-all politics: How Washington made the rich richer - and turned its back on the middle class*. Simon and Schuster.
- Hall, Richard L and Alan V Deardorff. 2006. "Lobbying as Legislative Subsidy." *American Political Science Review* 100(1):69–84.
- Hill, Kim Quaile and Patricia A Hurley. 2002. "Symbolic Speeches in the U.S. Senate and Their Representational Implications." *The Journal of Politics* 64(1):219–231.
- Hollibaugh, Gary E., Lawrence S. Rothenberg and Kristin K. Rulison. 2013. "Does It Really Hurt to Be Out of Step?" *Political Research Quarterly* 66(4):856–867.
- Hwang, June, Kosuke Imai and Alex Tarr. 2019. "Automated Coding of Political Campaign Advertisement Videos: An Empirical Validation Study."

- Jones, David R. 2003. "Position Taking and Position Avoidance in the U.S. Senate." *The Journal of Politics* 65(3):851–863.
- Jurafsky, Daniel and James H. Martin. 2008. *Speech and Language Processing: An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition*. Upper Saddle River: Pearson/Prentice Hall.
- Kalla, Joshua and Ethan Porter. 2019. "Correcting Bias in Perceptions of Public Opinion Among American Elected Officials: Results from Two Field Experiments." *OSF Preprints* .
- Kanthak, Kristin and Jonathan Woon. 2015. "Women Don't Run? Election Aversion and Candidate Entry." *American Journal of Political Science* 59(3):595–612.
- Kazin, Michael. 1998. *The Populist Persuasion*. Vol. 8.
- Key, V.O. 1949. *Southern Politics in State and Nation*. New York: Alfred A. Knopf.
- Kiesling, Scott Fabius. 1998. "Men's identities and sociolinguistic variation: The case of fraternity men." *Journal of Sociolinguistics* 2(1):69–99.
- Klofstad, Casey A., Rindy C. Anderson and Susan Peters. 2012. "Sounds like a winner: Voice pitch influences perception of leadership capacity in both men and women." *Proceedings of the Royal Society B: Biological Sciences* 279(1738):2698–2704.
- Knox, Dean and Christopher Lucas. 2019. "A Dynamic Model of Speech for the Social Sciences."
- Kostadinova, Tatiana and Barry Levitt. 2014. "Toward a theory of personalist parties: Concept formation and theory building." *Politics and Policy* 42(4):490–512.
- Ladam, Christina and Jeffrey J Harden. 2017. "Prominent Role Models : High-Profile Female Politicians and the Emergence of Women as Candidates for Public Office."
- Laustsen, Lasse and Alexander Bor. 2017. "The relative weight of character traits in political candidate evaluations: Warmth is more important than competence, leadership and integrity." *Electoral Studies* 49:96–107.
- Lawless, J. L. 2004. "Politics of Presence? Congresswomen and Symbolic Representation." *Political Research Quarterly* 57(1):81–99.
- Lawless, Jennifer L and Richard L Fox. 2004. "Why Don't Women Run for Office ?" *Brown Policy Report* .

- Lee, Frances E. 2016. *Insecure Majorities: Congress and the Perpetual Campaign*. Chicago: University of Chicago Press.
- Leibovich, Mark. 2014. *This Town*. New York: Penguin Group.
- Lenz, Gabriel S. 2009. "Learning and opinion change, not priming: Reconsidering the priming hypothesis." *American Journal of Political Science* 53(4):821–837.
- Lieberman, Mark. 2008. "Empathetic -in' October."
 URL: <http://languagelog.ldc.upenn.edu/nll/?p=732>
- Lieberman, Mark. 2011. "Pawlenty's linguistic 'southern strategy'?"
 URL: <http://languagelog.ldc.upenn.edu/nll/?p=3032>
- Lipinski, Daniel. 2004. *Congressional Communication: Content & Consequences*. Ann Arbor: University of Michigan Press.
- Maestas, Cherie. 2003. "The Incentive to Listen: Progressive Ambition, Resources, and Opinion Monitoring among State Legislators." *The Journal of Politics* 65(2):439–456.
- Maltzman, Forrest and Lee Sigelman. 1996. "The Politics of Talk: Unconstrained Floor Time in the U. S. House of Representatives." *The Journal of Politics* 58(3):819–830.
- Mansbridge, Jane. 1999. "Should Blacks Represent Blacks and Women Represent Women? A Contingent 'Yes'." *The Journal of Politics* 61(03):628.
- Mayhew, David. 1974. *Congress: The Electoral Connection*. Yale University Press.
- McGraw, Kathleen M., Richard Timpone and Gabor Bruck. 1993. "Justifying controversial political decisions: Home style in the laboratory." *Political Behavior* 15(3):289–308.
- Miller, Warren E. and Donald E. Stokes. 1963. "Constituency Influence in Congress." *The American Political Science Review* 57(1):45–56.
- Moffitt, Benjamin. 2016. *The Global Rise of Populism: Performance, Political Style, and Representation*. Stanford University Press.
- Moore, Emma and Robert Podesva. 2009. "Style, indexicality, and the social meaning of tag questions." *Language in Society* 38(4):447–485.
- Mutz, Diana C. and Byron Reeves. 2005. "The New Videomalaise: Effects of Televised Incivility on Political Trust." *American Political Science Review* 99(1):1–15.

- Niven, David and Jeremy Zilber. 2001. "Do women and men in congress cultivate different images? Evidence from congressional web sites." *Political Communication* 18(4):395–405.
- Nunberg, Geoff. 2008. "Palin's tactical g-lessness." .
URL: <http://languagelog.ldc.upenn.edu/nll/?p=733>
- Page, Benjamin I. 1976. "The Theory of Political Ambiguity." *The American Political Science Review* 70(3):742–752.
- Page, Benjamin I., Larry M. Bartels and Jason Seawright. 2013. "Democracy and the policy preferences of wealthy Americans." *Perspectives on Politics* 11(1):51–73.
- Pearson, Kathryn and Eric McGhee. 2013. "Should Women Win More Often than Men? The Roots of Electoral Success and Gender Bias in U.S. House Elections." *SSRN* .
- Potter, Claire Bond. 2019. "Men Invented 'Likability.' Guess Who Benefits." .
URL: <https://www.nytimes.com/2019/05/04/opinion/sunday/likeable-elizabeth-warren-2020.html>
- Rheault, Ludovic and Sophie Borwein. 2019. "Multimodal Techniques for the Study of Affect in Political Videos." .
- Ricks, Jacob I. 2018. "The Effect of Language on Political Appeal: Results from a Survey Experiment in Thailand." *Political Behavior* pp. 1–22.
- Rogers, Todd and David W. Nickerson. 2013. "Can Inaccurate Beliefs About Incumbents be Changed? And Can Reframing Change Votes?" *SSRN Electronic Journal* .
- Rothenberg, Lawrence S. and Mitchell S. Sanders. 2000. "Severing the Electoral Connection: Shirking in the Contemporary Congress." *American Journal of Political Science* 44(2):316–325.
- Sanbonmatsu, Kira. 2002. "Political Parties and the Recruitment of Women to State Legislatures." *The Journal of Politics* 64(3):791–809.
- Sandoval, Steven, Visar Berisha, Rene L. Utianski, Julie M. Liss and Andreas Spanias. 2013. "Automatic assessment of vowel space area." *The Journal of the Acoustical Society of America* 134(5):EL477–EL483.
- Saward, Michael. 2014. "Shape-shifting representation." *American Political Science Review* 108(4):723–736.

- Schattschneider, E.E. 1960. The Scope and Bias of the Pressure System. In *The Semi-Sovereign People: A Realist's View of Democracy in America*, ed. E.E. Schattschneider. Hynsdale, NY: The Dryden Press.
- Scherer, Nancy and Brett Curry. 2010. "Does descriptive race representation enhance institutional legitimacy? the case of the U.S. courts." *Journal of Politics* 72(1):90–104.
- Schilling-Estes, Natalie. 2003. Investigating Stylistic Variation. In *The Handbook of Language Variation and Change*, ed. J. K. Chambers, Peter Trudgill and Natalie Schilling-Estes. John Wiley & Sons pp. 549–572.
- Schlozman, Kay L., Sidney Verba and Henry E. Brady. 2012. *The Unheavenly Chorus. Unequal Political Voice and the Broken Promise of American Democracy*. Princeton, NJ: Princeton University Press.
- Shepsle, Kenneth. A. 1972. "The Strategy of Ambiguity : Uncertainty and Electoral Competition." *The American Political Science Review* 66(2):555–568.
- Spierings, Niels and Andrej Zaslove. 2017. "Gender, populist attitudes, and voting: explaining the gender gap in voting for populist radical right and populist radical left parties." *West European Politics* 40(4):821–847.
URL: <http://dx.doi.org/10.1080/01402382.2017.1287448>
- Story, Brad H. and Kate Bunton. 2017. "Vowel space density as an indicator of speech performance." *The Journal of the Acoustical Society of America* 141(5):EL458–EL464.
- Stuart-Smith, Jane, Claire Timmins and Fiona Tweedie. 2007. "'Talkin' Jockney'? Variation and change in Glaswegian accent." *Journal of Sociolinguistics* 11(2):221–260.
- Thai, Xuan and Ted Barrett. 2007. "Biden's description of Obama draws scrutiny."
URL: <http://www.cnn.com/2007/POLITICS/01/31/biden.obama/>
- Tomz, Michael and Robert P. Van Houweling. 2009. "The electoral implications of candidate ambiguity." *American Political Science Review* 103(1):83–98.
- Whitfield, Jason A. and Alexander M. Goberman. 2014. "Articulatory-acoustic vowel space: Application to clear speech in individuals with Parkinson's disease." *Journal of Communication Disorders* 51:19–28.
- Yiannakis, Diana Evans. 1982. "House Members' Communication Styles: Newsletters and Press Releases." *The Journal of Politics* 44(4):1049–1071.

Yuan, Jiahong and Mark Liberman. 2011. "Automatic detection of "g-dropping" in American English using forced alignment." *2011 IEEE Workshop on Automatic Speech Recognition and Understanding, ASRU 2011, Proceedings* pp. 490–493.

Appendix 1: Technical Appendix

This section of the appendix attempts to explain commonly used techniques from natural language processing, mechanical/electrical engineering and phonetics for a political science audience. As such, it cannot cover every aspect in full detail. The reader is encouraged to refer to the relevant literature, such as Jurafsky and Martin (2008), for further reference.

Audio Data

At a fundamental, physical level, sound is a wave traveling through and disturbing a medium, such as the air. When undisturbed, the medium is in equilibrium. When a wave propagates through it, the **sound pressure** causes a corresponding disturbance, also referred to as **amplitude**. The high points (i.e. maximum distance to the equilibrium) of the wave are referred to as **crests** and the low points as **troughs**. When traveling through air, a sound wave moves at a speed of about 343m/s (depending on the temperature). The **frequency** f of a wave describes the number of cycles it undergoes per time period T , usually per second - which is then described as hertz, or Hz.

$$f = \frac{1}{T}$$

This means that for a wave traveling at 20Hz, 0.05s pass between two crests (i.e. a **wavelength**).

The method used here for representing an analog signal in digital form is pulse code modulation. To do so, the amplitude of the signal is sampled at regular intervals. The number of these intervals within a given time frame (i.e. the frequency) corresponds to the **sampling rate**. Each cycle of a wave needs to be represented by at least two data points, one for the positive and one for the negative section of the wave. Sound perceptible to humans occurs between 20 Hz and 20,000 Hz. Hence, a sampling rate of 40,000 Hz (i.e. 40,000 data points per second) would be necessary to adequately represent this information. The audio tracks for most of the videos used in this analysis have a sampling rate of 44,100 Hz. Human speech generally only occurs at frequencies below 10,000 Hz, and the formants (see below) which constitute my most important form of data occur below 5,000 Hz for males and 5,500 Hz for females (young children can go up to 8,000 Hz, but this is of no concern to this analysis). Consequently, encoding audio at such high quality is not strictly necessary for my purposes.

Furthermore, the **bit depth** determines the number of possible values at each interval.

For example, a sampling rate of 16000 Hz means there are 16000 samples per second, and a bit depth of 16bit means that each sample has a resolution of 2^{16} possible values. This means that integers between -32768 and 32767 can be represented, entailing a fairly high level granularity.

Another concept of some relevance here is the number of channels. Most audio, including the videos scraped from YouTube in this paper, is in stereo format, meaning two channels. For the analyses in this paper, this is entirely irrelevant, so I convert all of my audio data to mono.

Hence, both sampling rate and bit depth determine the audio quality. These concepts, along with the number of channels, determine the bitrate as following:

$$\text{Bit rate} = \text{sampling rate} * \text{bit depth} * \text{channels}$$

As an example, uncompressed .wav files usually store data at 16-bit, 44.1 kHz and consist of two audio channels (i.e. stereo), so 1 hour of audio requires 635.04 MB of storage (about the size of a CD). Even a compressed¹¹ MP3 file with a bit rate of 128kB/s (i.e. the number of bits used for each second of audio) still weights in at 57.6 MB. It follows that the storage demands for this project are rather large. Since, as discussed above, human speech does not fill the entire spectrum of human aural perception, I address this problem by using a lower sampling rate: In this paper, I use .wav files with a sampling rate of 16000 Hz, a bit depth of 16bit, and one channel, resulting in a bitrate of 256Kbit/s.

The Fourier transform

The Fourier transform is an extremely commonly used technique in the processing and analysis of audio data. In this paper, it is used at several steps, such as the extraction of formants, speech diarization and speaker recognition. It is used to convert a signal from the time to the frequency domain.

There are two types of the Fourier transform - one for continuous and one for discrete data. Since audio data in the digital domain is in discrete form, I generally rely on the latter.

Continuous Fourier transform. Applied to an analog signal of potentially infinite length.

$$X(F) = \int_{-\infty}^{\infty} x(t)e^{-j2\pi Ft} dt \quad (1)$$

Discrete Fourier transform. The discrete Fourier transform is applied to a windowed

¹¹Audio compression is based on filtering out information that is imperceptible to the human ear and is therefore redundant.

portion $x[n]...x[m]$ of a signal.

$$X_k = \sum_{n=0}^{N-1} x_n e^{-j2\pi kn/N} \quad (2)$$

The output X_k , represents the k -th frequency bin. Since the audio signal is sampled at discrete points, this means that an amplitude is being calculated at each frequency. There are n such frequency bins. For example, if a signal is sampled 8 times, n will be $n = [0, 1, 2, ..., 7]$. The points k at which the frequencies are sampled are all multiples k of the fundamental frequency $\frac{2\pi}{T}$ in the continuous case, and $\frac{2\pi}{N}$ in the discrete case.

Mel-frequency cepstral coefficients (MFCCs)

Mel-frequency cepstral coefficients (MFCCs) are a commonly used feature in machine learning applied to audio tasks. MFCCs denominate how much energy exists in regions of the frequency domain. This matters because human hearing cannot perceive very small differences in frequencies, but is generally better at doing so for lower frequencies. Consequently the filterbank determines the increasing distances of the ‘bins’ into which frequencies are grouped.

To this end, frequencies are converted to mels according to the following equation¹²:

$$M(f) = 1125 \ln(1 + \frac{f}{700}) \quad (3)$$

To convert back to frequency, the following equation is used:

$$f = 700(\exp \frac{m}{1125} - 1) \quad (4)$$

To calculate the MFCCs, the audio signal is portioned into a set of windows, each of which consists of a number of frames. Then, the discrete fourier transform, as outlined above, is applied to such a window. Then, the periodogram power spectral density estimate is calculated by squaring the absolute value of the output of the DFT and dividing by the number of samples in the window. Then, the mel-spaced filterbank is applied to this periodogram, yielding 26 coefficients. After taking the log of these numbers and applying the discrete cosine transform, I am left with 26 cepstral coefficients for each frame. In this paper, MFCCs are used for both speech diarization and speaker recognition. For the formant extraction, Linear Prediction Coefficients (LPCs), a similar type of feature, are used. The speaker recognition program also makes use of delta MFCCs, which denominate the change in MFCCs between frames. For further reference,

¹²Note that there is no ubiquitous equation for this. The constant in front of the log is not always exactly the same.

see Jurafsky and Martin (2008).

Speaker Identification with Gaussian Mixture Models (GMMs)

Gaussian mixture models are generative models describing the distribution of data. They are frequently used as a probabilistic clustering model, which addresses some of the shortcomings of k-means clustering that some political scientists might be familiar with. GMMs assume a multivariate Gaussian distribution, which gives them more flexibility than, for example, k-means clustering. The analogue of clusters in this model are its components. The data-generating process assumes the selection of a component, and the subsequent generation of a data point from a normal distribution according to the model parameters. These parameters are learned via Expectation Maximization (EM).

In the case of speaker identification, I use a GMM with 16 components and train one model for each of the senators in the sample. Essentially, the GMM learns, in an unsupervised manner, how the MFCCs for the speaker's training data were generated from a multivariate normal distribution. At test time, the model computes, for each frame and for each component, the log-likelihood that this test data was generated, given the model parameters of each speaker model. These log-likelihoods are then summed up, and the speaker model with the highest log-likelihoods is predicted to be the actual speaker of the sample.

Appendix 2: Tables & Figures

Table 1: Linear regression. The effect of setting (campaign/Congress) on vowel space area. The table shows that senators speak more colloquially in a campaign context.

	<i>Dependent variable:</i>			
	Vowel space area			
	(1)	(2)	(3)	(4)
Setting (Campaign)	−0.175*** (0.042)	−0.051*** (0.019)	−0.177*** (0.042)	−0.048** (0.019)
Gender (Male)	−0.281*** (0.063)	−0.231*** (0.061)	−0.281*** (0.064)	−0.229*** (0.062)
Video duration	0.010 (0.011)	0.011 (0.011)	0.009 (0.011)	0.010 (0.011)
Party (Republican)	−0.154*** (0.055)	−0.151*** (0.055)	−0.154*** (0.055)	−0.151*** (0.055)
South	0.010 (0.056)	0.007 (0.056)	0.021 (0.056)	0.018 (0.056)
% Rural population	0.134 (0.170)	0.145 (0.170)	0.067 (0.163)	0.076 (0.163)
Male*Campaign	0.157*** (0.047)		0.163*** (0.047)	
Constant	13.226*** (0.069)	13.180*** (0.067)	13.239*** (0.069)	13.191*** (0.067)
Random intercept				
Observations	3,466	3,466	3,429	3,429
Log Likelihood	−2,041.499	−2,044.944	−2,017.466	−2,021.379
Akaike Inf. Crit.	4,102.997	4,107.888	4,054.931	4,060.759
Bayesian Inf. Crit.	4,164.505	4,163.245	4,116.331	4,116.019

Note:

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

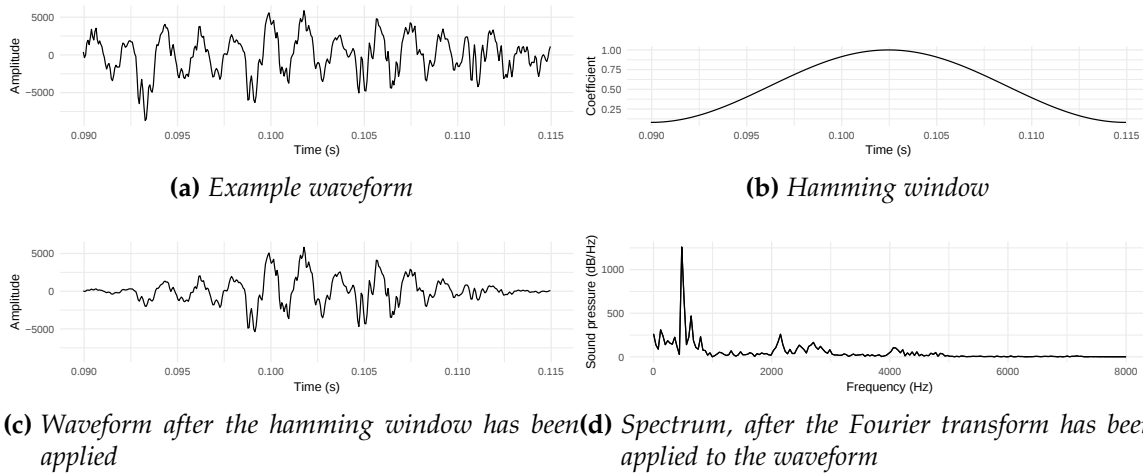


Figure 6: These four figures show the process of transforming a basic waveform into a spectrum. This spectrum will then, in the next step, be turned into a spectrogram, from which formants can be identified.

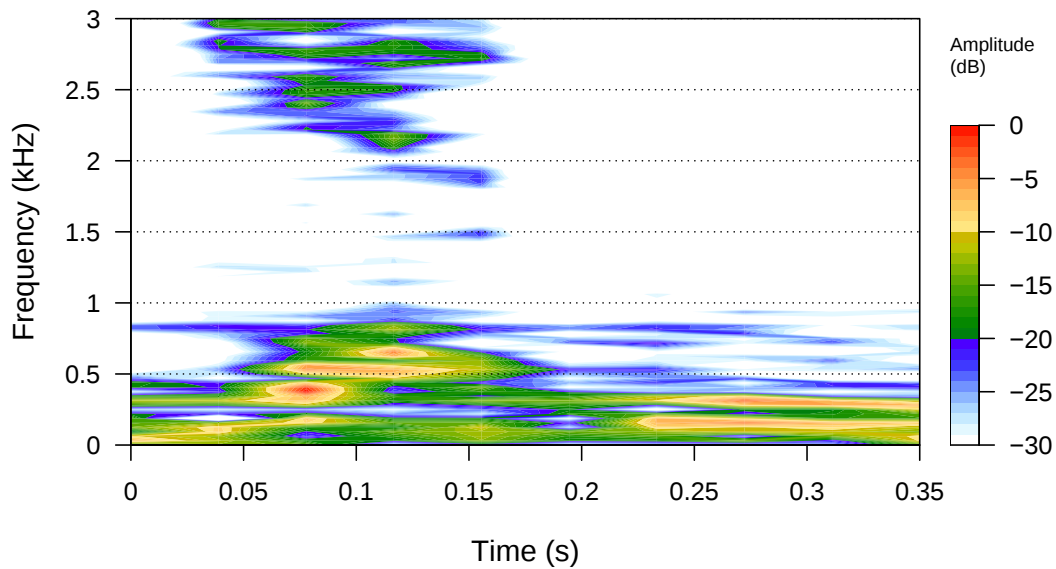


Figure 7: Spectrogram of a single word. The spectrogram is essentially a spectrum over time, plotted as a heatmap. The high-density areas that go on for some time correspond to the formants, from the bottom up. So F1 corresponds to the bottom “line”, F2 the next, and so on.