



ROCKWOOL Foundation Berlin

Institute for the Economy and the Future of Work (RFBerlin)

DISCUSSION PAPER SERIES

080/26

Do Elections Moderate or Polarize Political Rhetoric?

Tito Boeri, Nina Nikiforova, Guido Tabellini

Do Elections Moderate or Polarize Political Rhetoric?

Authors

Tito Boeri, Nina Nikiforova, Guido Tabellini

Reference

JEL Codes: P,H

Keywords: Electoral competition, Populism, Partisanship, Polarization

Recommended Citation: Tito Boeri, Nina Nikiforova, Guido Tabellini (2026): Do Elections Moderate or Polarize Political Rhetoric?. RFBerlin Discussion Paper No. 080/26

Access

Papers can be downloaded free of charge from the RFBerlin website: <https://www.rfberlin.com/discussion-papers>

Discussion Papers of RFBerlin are indexed on RePEc: <https://ideas.repec.org/s/crm/wpaper.html>

Disclaimer

Opinions and views expressed in this paper are those of the author(s) and not those of RFBerlin. Research disseminated in this discussion paper series may include views on policy, but RFBerlin takes no institutional policy positions. RFBerlin is an independent research institute.

RFBerlin Discussion Papers often represent preliminary or incomplete work and have not been peer-reviewed. Citation and use of research disseminated in this series should take into account the provisional nature of the work. Discussion papers are shared to encourage feedback and foster academic discussion.

All materials were provided by the authors, who are responsible for proper attribution and rights clearance. While every effort has been made to ensure proper attribution and accuracy, should any issues arise regarding authorship, citation, or rights, please contact RFBerlin to request a correction.

These materials may not be used for the development or training of artificial intelligence systems.

Imprint

RFBerlin
ROCKWOOL Foundation Berlin –
Institute for the Economy
and the Future of Work

Gormannstrasse 22, 10119 Berlin
Tel: +49 (0) 151 143 444 67
E-mail: info@rfberlin.com
Web: www.rfberlin.com



Do Elections Moderate or Polarize Political Rhetoric?*

Tito Boeri[†] Nina Nikiforova[‡] Guido Tabellini[§]

March 17, 2026

Abstract

We study the communication strategies on Twitter/X of 367 political leaders in 21 countries, focusing on electoral competition between populists and non-populists. We measure polarization by the ease with which the leader can be classified as populist or not, conditional on his tweet. We find that political rhetoric becomes more polarized before and around election dates. This happens because, in pre-electoral quarters, opposite leaders are more likely to: i) talk about different topics, and ii) frame differently the same issues. Our results are consistent with competing politicians targeting different voters, rather than appealing to the same swing voters.

Keywords: Electoral competition, Populism, Partisanship, Polarization.
JEL classification: P, H

*We thank Fondazione Rodolfo De Benedetti for access to Twitter data, Leonardo Bianchi, Ginevra Casini, Marco Cerundolo and Gabriele Nespola for outstanding research assistance, and Elliott Ash, Jim Fearon, Torsten Persson, Andrea Prat, Jesse Shapiro and seminar participants at Bocconi University, Cambridge, CEMFI, King's College, Oxford, the IOG meeting at Berkeley and the CEPR media conference at Bocconi for helpful comments. Tabellini gratefully acknowledges financial support from the Italian Ministry of University and Research, D.D. n. 1802 del 21/11/2024 – BANDO FIS 3 - progetto FIS-2024-05869 - CUP J53C25002480001

[†]Department of Economics and IGIER, Bocconi University; IZA; CEP-LSE

[‡]Department of Economics, Bocconi University

[§]Department of Economics and IGIER, Bocconi University; CEPR; CES-Ifo

1 Introduction

Do elections induce moderation or polarization between competing politicians? If competing candidates seek to persuade the same voters, as in standard theories of Downsian or probabilistic voting, then most theories of electoral competition predict policy convergence (Persson and Tabellini, 2002). In this approach, divergence is due to intrinsic partisan ideologies, that make candidates accept vote losses to be closer to their preferred positions (Alesina, 1988 and Besley and Coate, 1997). An alternative view is that elections have a polarizing effect. As discussed below, this can happen in a variety of situations, but the most compelling one is that candidates target different constituencies. If so, elections create incentives to diverge, as each candidate seeks to please or mobilize its targeted voters (Glaeser, Ponzetto and Shapiro (2005), Gennaioli and Tabellini (2025) and Prummer (2020)). Despite a large empirical literature discussed below, the evidence on which view is closer to the truth is still inconclusive.¹

To shed light on this question, we study how the communication strategies of politicians on Twitter/X are affected by the imminence of national (legislative and presidential) elections. We ask whether political rhetoric becomes more or less polarized between opposite groups of politicians before or after the elections, compared to non-election periods.

We study about 3.4 million tweets sent by 367 political leaders in 21 Western democracies during 2013-2022, and exploit staggered election dates across countries. Messages on Twitter/X, unlike electoral programs or speeches, provide high frequency data, which allows us to study time variation as the election approaches. Tweets are more followed than party manifestos or speeches, particularly by journalists and people interested in politics, they offer more freedom of expression than traditional media, and they provide a free platform for politicians challenging the establishment, as is the case of most populist leaders. For these reasons, platforms like Twitter/X have been a key arena of political competition, and impact on the salience of topics such as crime, illegal

¹Note that the key difference between these two opposing views is not about the relevant voting margin: whom to vote for vs. whether to vote. The key issue is who hears the campaign messages. Even when turnout is the relevant margin, vote maximizing candidates face incentives to diverge only if, when mobilizing their core voters, they do not also mobilize those who would vote for the opponent. And vice versa, even if turnout is fixed, candidates may find it optimal to diverge if their messages reach different groups of voters.

migration or mass shootings, as well as on how these topics are discussed (Battiston, 2022 and Zhang et al., 2025). Finally, compared to party manifestos or political speeches and to other media, tweets can be more easily targeted to specific audiences, and politicians can monitor the feedback received and by what type of voter. As discussed below, this can contribute to explain why our findings differ from those in other similar studies that have focused on candidate websites and manifestos.

Our units of observation are the bigrams (i.e. two word sequences) used by each politician during a quarter. The main challenge is that the sample of observed bigrams is small, relative to the very large number of bigrams that could be used. As explained by Gentzkow, Shapiro and Taddy (2019) (henceforth GST), simple measures related to word counts over-estimate deliberate rhetorical polarization, because many bigrams are randomly chosen. We thus follow their pioneering approach, and use a Lasso regression method to correct for finite sample bias. As in GST, we mostly measure rhetorical polarization by *partisanship*, namely the ease with which one can correctly predict a politician’s type, based on his words. For each politician, partisanship provides a quarterly measure of how polarized their tweets are, relative to politicians with the same observed characteristics but of the opposite political type.

We estimate different measures of partisanship. Our main focus is on populist vs non-populist politicians. But we also consider left vs right, as well as a partition between four political types: left-wing populist, right-wing populist, and similarly for non-populists. These classifications are defined ex-ante, based on ChatGPT and party affiliations. The underlying assumption is that the political type influences word usage in similar ways across languages and countries, once we control for other observed individual characteristics. We also estimate a quarterly Chi-square measure of within-country heterogeneity of political speech, correcting for finite sample bias. This Chi-square measure is closely related to partisanship, but does not entail any assumption about what are the relevant political types. We then use these measures of rhetorical polarization as dependent variables in an event study, where we compare polarization in a window of 7 quarters around national elections with quarters outside of this window.

There are several reasons for focusing on populist vs non-populist competition. This rivalry is arguably the main dimension of political competition in most Western democracies in this period. Populist leaders are or have been in

government in the US, Brazil and Mexico, they are part of governing coalitions in the Netherlands, Austria, Italy and Sweden, influence politics from the outside in France, and command large majorities in Eastern German states. The rise of populist challengers is the other side of the coin of the decline of mainstream parties. Moreover, populist movements tend to be more similar and more connected across countries, compared to other political groups. Populist politicians in our sample have more international ties, measured by Twitter links between political leaders, compared to non-populist leaders. The populist / non-populist classification is also very stable over time (only one leader, Pawel Piotr Kukiz from Poland, appears to have switched from non-populist to populist during our sample period, according to ChatGPT, and there are 3 more politicians for whom the evidence of switching is mixed).

Unlike most of the literature on populism, we focus on leaders rather than parties. This seems appropriate when studying the supply of populism, which is often based on the emergence of a charismatic leader.

Our main result is that political rhetoric becomes more polarized two quarters before the election and remains so up to and including the election quarter, relative to its average level in quarters more distant from the election, while in post-election quarters it falls back to the average. This finding is very robust and has a large magnitude for populist vs non-populist politicians. On average during the election quarter, the predictability of populist vs non-populist types exceeds the predictability of the same politician in quarters distant from the election by about 10-15% of a standard deviation of predictability during the entire sample. This magnitude is comparable to the peak effects of an aggregate shock that polarized populist vs non-populist speech, such as the 2015 refugee crisis in Europe.

These patterns are stronger in countries under plurality rules and presidential regimes, and in presidential elections, where typically competition is between two major parties, suggesting that they are not due to the presence of a large number of candidates.

Results are qualitatively similar, but smaller, for the Chi-square measure of polarization that does not rely on any prior classification of leaders into political types. During the election quarter the Chi-square measure rises on average by about 5% of a standard deviation, relative to quarters distant from the election. Here too, polarization rises two quarters before the election, and after the election it falls back to its average level. On the other hand, for left

vs right political types we find much less rhetorical polarization on average, and no evidence of significant pre-electoral cycles.

Prior research has characterized populism as a "thin-centered ideology", meaning that it places relatively little emphasis on detailed policy content beyond a general aversion to immigration and globalization. This is confirmed by our analysis. We classify tweets by their topic, for 20 policy and non-policy topics, and we then assign bigrams to each topic based on the frequency with which they appear in tweets on that topic (about 40% of bigrams are not assigned to any topic). We then estimate the posterior that the politician is populist, conditional on the bigram's topic. Topics for which the average posterior over the entire sample is above (below) $1/2$ are typical of populist (non-populist) politicians. As expected, non-populist politicians are more likely to tweet about policy issues, with the exception of immigration and, after COVID, health policy, which instead are distinctively populist (particularly immigration). On the other hand, populists are more likely to tweet against the establishment, engage in self-promotion, or to mention specific parties, politicians and elections.

These differences between populist and non-populist politicians become more pronounced before and during the election. Increased rhetorical polarization two quarters before and up to the election quarter is observed *within topics* (i.e. the same issues are framed differently), but particularly *between topics* (i.e. populist vs non-populist politicians are more likely to tweet about different topics). Specifically, as the election approaches, observing a bigram belonging to a non-populist topic (i.e. mainly a policy topic other than immigration and post COVID health) becomes more informative that the speaker is non-populist.

Although we cannot tell who is responsible for this increased rhetorical divergence before the election, whether populists or non-populists, we close by asking whether language becomes more or less populist as the election approaches. Following GST, we measure how distinctive each bigram is of populist speech over the entire sample period. We then ask whether populist bigrams are used more or less frequently as the election approaches. We find that on average tweets become more populist one quarter before the election and in the election quarter, relative to quarters distant from elections. This is due to very non-populist bigrams (those in the lower 20% of the distribution of populist distinctiveness) being used less often, by both political types. Many of

these bigrams refer to specific institutions or other countries, with which voters do not easily connect. In other words, closer to the election, all politicians are less likely to adopt a formal and cold style of communication, which is distinctive of non-populist rhetoric. Of course, this does not contradict the main finding of increased rhetorical polarization before the election, because partisanship (our measure of polarization) and populist distinctiveness are different concepts and they are measured differently.

Overall, these findings suggest that elections have a polarizing effect on the rhetoric of populist vs non-populist political leaders on Twitter/X. The results using the Chi-square measure, which is agnostic about political types, suggest that this applies to all political leaders, although the effect is stronger between populists vs non-populists.

This is unlikely to be due to changes in the audience: the number of politicians' followers in Twitter/X does not fluctuate around election dates (Van Kessel et al., 2020). A more plausible interpretation, instead, is that electoral incentives exert a stronger influence on political communication in pre-election periods, compared to periods after or distant from elections, and these incentives induce competing politicians to diverge. During an electoral campaign, communication is focused on the short run goal of persuading and mobilizing voters. In off-election periods, instead, politicians can afford to let their Tweets be shaped by their ideology and personal traits, or to react to external events. Although there are several possible reasons why electoral incentives may induce polarization, a prominent one is that populist and non-populist leaders target different groups of voters, and they seek to mobilize their constituencies before the elections. Indeed, the more precise targeting of core supporters afforded by Twitter can explain why our results differ from those on other forms of political communication.²

We contribute to several lines of research. First, we apply GST's methodology to study different data and a novel question - GST did not study the effect of elections on political rhetoric. Polarization on Twitter between populist and non-populist politicians is larger and more volatile than in congressional speeches of US Republicans and Democrats. We also extend their measure of polarization to more than two groups of politicians and relate it to the

²Until 2023 Twitter analytics allowed to recover breakdowns of followers by gender, location, interests (with categories such as "conservative politics"), and device usage. Thus, in addition to relying on self-selection of followers, leaders could assess whether their communication campaigns were effective in targeting specific audiences.

Chi-square measure of heterogeneity.³

Second, we contribute to the literature on whether elections induce political convergence or divergence. Besides the research discussed above, on swing voters vs endogenous turnout and on voters' informational asymmetries, other theoretical papers have shown that elections create incentives to diverge, if voters have convex policy preferences (Kamada and Kojima, 2014), or candidates are risk averse and care about their vote share (Bonomi, 2024), or districts are heterogeneous and there is within party disagreement (Polborn and Snyder, 2017). Empirical research on these issues has not reached conclusive results. In line with our divergence findings, Ash, Morelli and Van Weelden (2017) show that floor speeches of US senators devote more time to divisive issues when they are up for election. On the other hand, using data of candidate websites in US House elections and manifestos for French elections, Di Tella et al. (2025) find evidence of convergence in ideology and rhetorical complexity between primaries / first round elections and final elections. The difference with our results could be due to Twitter enabling better targeting of voters, compared to official candidate or party positions, or to primaries/first round elections inducing even more divergence than final elections. This second interpretation is in line with the findings by Brady, Han and Pope (2007) and Fowler and Fu (Forthcoming).

Our results also speak to a related literature in political science, which has evaluated empirically the tradeoff in vote shares between converging to the center to capture swing voters, vs diverging to mobilize more extreme voters. Here too, results are inconclusive. An earlier literature suggests that pre-electoral divergence in Congressional roll call votes is punished at the election (Canes-Wrone, Brady and Cogan, 2002). Similarly, Hall (2015) and Hall and Snyder Jr (2015) find that, when a more extremist candidate wins the primaries, he is penalized in the general election. On the other hand, Bonica, Rhee and Studen (2025) argue that these findings are confounded by endogenous turnout and ballot composition effects. When these are taken into account, they find that the small electoral penalty for divergence is swamped by the large electoral gains of mobilizing core voters. Hill (2017) reaches similar conclusions with a different methodology.

Third our paper contributes to the literature on populism and social me-

³Fiva, Nedregård and Øien (forthcoming) apply GST's methods to study different dimensions of heterogeneity in legislative speeches of Norwegian MPs.

dia. Most of it seeks to explain voters' behavior - cf. [Gurieva and Papaioannou \(2022\)](#) and [Gurieva, Melnikov and Zhuravskaya \(2021\)](#) and [Manacorda, Tabellini and Tesei \(2022\)](#). But the supply side is equally important, and much less is known about it, particularly on the communication strategies of populists in proximity of elections, and on how non-populist politicians react to these strategies. The finding that polarization increases ahead of elections is at odds with the view that social media enable politicians to wage a permanent campaign ([Gibson, 2020](#)), "making it difficult to distinguish political communication in non-electoral periods from that in electoral periods" ([Dommett et al., 2025](#)). Our results suggest that communication is significantly different over the electoral cycle.⁴

Last, but not least, we contribute to research on political polarization. A consensus is emerging that politicians' messages and propaganda influence voters' beliefs and fuel polarization among citizens (see [Zhang et al. \(2025\)](#) and [Klein \(2020\)](#)). Nevertheless, it is less clear why politicians behave in this way, whether for opportunistic reasons or due to their ideological convictions. Our results are consistent with the predictions of [Gennaioli and Tabellini \(2025\)](#) and [Bonomi \(2024\)](#), who show that opportunistic politicians may find it optimal to polarize the electorate.

The paper is structured as follows. Section 2 describes the data. Section 3 describes and extends the methodology developed by GST. Section 4 presents some aggregate patterns in the data and shows that our measure of partisanship reacts to aggregate shocks like the refugee crisis or Covid. Section 5 presents our results on partisanship over the electoral cycle, while Section 6 evaluates the underlying mechanisms and discusses robustness. Section 7 assesses whether rhetoric on average becomes more or less populist as the election approaches. Section 8 concludes.

2 Data

2.1 Tweets and Politicians

Data on tweets of politicians were collected through Twitter Academic API, version 2. The sample covers the period 2013-22 in 21 countries selected

⁴Research in political science has studied populist communication strategies, but without a focus on the electoral cycle - cf. [Cassell \(2023\)](#) and [Aslanidis \(2018\)](#).

also based on the prominence of populist political leaders (Australia, Austria, Belgium, Brazil, Czech Republic, Denmark, Finland, France, Germany, Hungary, Italy, Mexico, Netherlands, Norway, Poland, Portugal, Slovenia, Spain, Sweden, United Kingdom, United States).⁵

Politicians were selected based on their leadership positions, according to the following criteria: (i) All candidates for Prime Minister/President in general elections between 2001 and 2022. (ii) Leaders of main political parties with vote share above 5% in at least one general election between 2010 and 2022. (iii) Influential political leaders at the local level or leaders of niche parties, even if the national vote share of their party is below 5 % (e.g., Nicola Sturgeon in the UK, Basque and Catalan political leaders in Spain, Giorgia Meloni in Italy). These criteria identified 496 leaders, of whom 367 had an active Twitter account, for a total of about 3.4 million tweets written between January 2013 and June 2022. In addition to the text of each tweet, we have the exact date and time in which they were issued.

A distinguishing feature of our paper is that (unlike most of the literature on populism) we focus on political leaders rather than on parties. We classify politicians along two dimensions: populist (P) / non-populist (NP) and left (L) / right (R) divides.

2.1.1 Populist vs. Non-Populist

To identify populist leaders, we draw on the classification of [Funke, Schularick and Trebesch \(2023\)](#) for those few leaders in their dataset who were active during the period covered by our analysis. For the remaining politicians, we conducted a manual classification informed by GPT-4o mini-generated responses to the following question: *“Is leader X from country Y commonly considered a populist leader?”*. A human coder then interpreted the lengthy and detailed answer produced by ChatGPT and on the basis of the explanation classified the leader as: populist, non-populist or ambiguous. The ambiguous category includes genuinely unclear cases, situations where ChatGPT was uncertain, or cases where leaders appeared to shift over time between populist and non-populist positions.⁶

⁵We thank the Rodolfo De Benedetti Foundation for giving us access to these data that they had collected.

⁶We also validated our classification against the fully automated classification produced by ChatGPT. There are only nine cases of complete disagreement between the two approaches – that is, politicians whom we classify as populists but whom the automated

After this step, 29 politicians were classified as *ambiguous*. We classify 16 of these as populist because their party is considered populist by at least one of Norris (2019), Rooduijn (2019) and Guriev and Papaioannou (2022). For the remaining 13 politicians we check that our results are robust to classifying them either way (cf. Section 6.3), and in our main analysis we include them among populist types.

Examples of leaders belonging to this residual group are Lula da Silva in Brasil, Nick Xenophon in Australia, and Marjan Sarec in Slovenia among others. Importantly, there are at most four politicians that, according to ChatGPT, may have switched from NP to P or vice versa in the period covered by our data. For these four politicians, we retain throughout the most recent classification. Appendix F describes how we established that. Appendix I lists all politicians and their political type. In total, we classify 86 politicians as populist and 281 as non-populist.

Our leader-based classification of populism is quite different from a party-based classification, and adds relevant information to it. About 86 % of populist leaders according to our definition operate in parties classified as populist in the Global Party Survey (Norris, 2020). However, there are also several politicians that we classify as non-populist although they belong to populist parties, such as Mitt Romney, Mike Pence and George Bush in the US, Gordon Brown, David Cameron, and Theresa May in the UK. Instead, Alexandria Ocasio-Cortez in the US and Mark Latham in Australia are examples of populist leaders in parties that are classified as non-populist.

Figure 1 illustrates the number of active politicians in each quarter disaggregated by their type, as well as the number of tweets they have sent. Politicians of both types are always active in all quarters (the minimum number of active populist politicians is always above 40).

Table 1 provides information on the characteristics of the two sets of politicians. Populist and non-populist leaders have approximately the same age, 3 out of 4 are males, and, on average, they began tweeting in 2012. What distinguishes P and NP leaders is their experience (P entered the national legislature at a later stage) and exposure to Government positions (higher for NP leaders).⁷ Some 80 % of NP leaders belong to mainstream parties, defined as

ChatGPT classification identifies as non-populists. Further details on this comparison are provided in Appendix F.

⁷The variable years in Govt measures the number of years in which the leader was either in a party in Government, or supporting the Government, or was herself a Minister

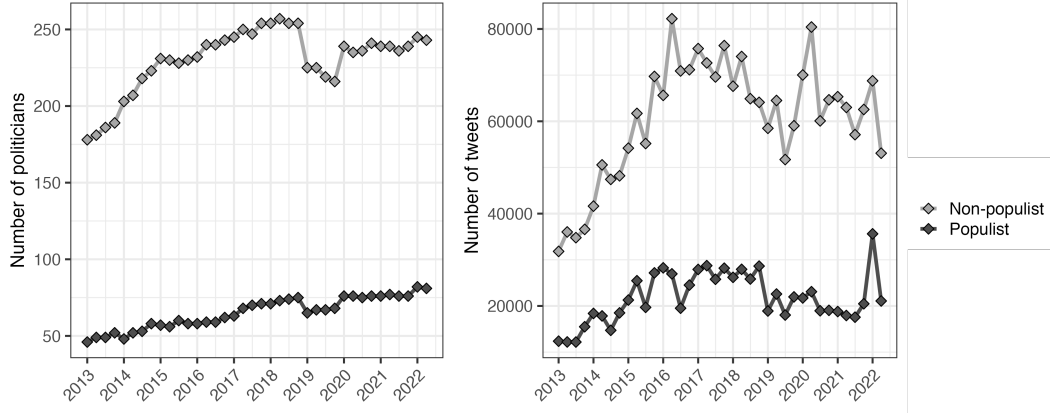


Figure 1: Number of politicians active on Twitter per calendar quarter, and number of tweets

parties that were in Government (or lent support to a Government) between 1990 and 2010.⁸ NP leaders also display, on average, higher educational attainments than P leaders, but the difference is not statistically significant at conventional levels. Almost 50 % of populist leaders belong to niche parties, based on a GPT-4o model classification that defines as niche parties those emphasizing a limited set of issues that do not coincide with the predominant economic Left/Right division (Meguid, 2005) (see Appendix G for the ChatGPT prompt). Differences in Government and legislative experience as well as in affiliation to either mainstream or niche parties are statistically significant at conventional levels. Overall, compared to NP leaders, P leaders have mostly the characteristics of *challengers*, because either they have not been in Government or they belong to parties representing issues not typically dealt with by the establishment.

2.1.2 Left vs Right

We relied on GPT-4o mini model to classify leaders along the L/R divide. We performed 4 iterations with the prompt displayed in Appendix H. Politicians belong to the "Right" group if GPT always unambiguously classified the politician in this category, and the same with "Left". For 82 politicians, one of the 4 GPT iterations said that the classification is ambiguous. To classify these 82 ambiguous cases, we relied on their party affiliations, using the most recent RILE index of the Manifesto project, which measures how much a party

⁸A leader is defined as mainstream even without government experience if she belonged to a party for a short period of time and that party was in government in a different period.

Table 1: Differences between Populists and Non-Populists, and Right and Left (means; difference with p-value in parentheses).

Variable	Populist vs Non-Populist			Left vs Right		
	P	NP	$\Delta(p)$	L	R	$\Delta(p)$
Age (in 2010)	44.43	48.32	-3.89 (0.01)	48.15	46.31	-1.85 (0.13)
Male	0.73	0.76	-0.03 (0.59)	0.74	0.78	0.05 (0.28)
Higher Education	0.86	0.91	-0.05 (0.26)	0.89	0.91	0.02 (0.50)
Year of First Tweet	2013.0	2012.4	0.58 (0.08)	2012.4	2012.7	0.39 (0.14)
Years in Govt	6.44	9.79	-3.35 (0.00)	9.28	8.56	-0.72 (0.48)
Legislative Experience	0.26	0.39	-0.14 (0.00)	0.34	0.40	0.07 (0.04)
Mainstream	0.58	0.82	-0.24 (0.00)	0.72	0.83	0.10 (0.02)
Niche Party	0.43	0.17	0.26 (0.00)	0.20	0.28	0.08 (0.07)

Notes. *Legislative Experience:* years passed in 2010 since they first entered the national legislature. *Mainstream:* equals 1 if the leader has at any point in their career belonged to a party which was in government or supported the government coalition in the years 1990-2010. *Niche parties* classification is produced by GPT-4o, and based on the definition of ‘niche’ proposed in Meguid (2005). See Appendix G for more details.

covers left-wing or right-wing issues.⁹ Based on this, we assigned 64 out of 82 leaders to the left and the remaining 18 to the right. We rely on GPT rather than party affiliations in order to capture within-party heterogeneity, like for the populist classification. Appendix H discusses the differences between our classification and the one that would result from only using party affiliations. Appendix I lists all politicians along the L-R divide.

⁹There are no iterations when GPT-4o mini model classification switched from classifying a politician from right to left or vice-versa, but just observations where the classification switched between right and ambiguous or between left and ambiguous or vice versa. The source for the manifesto data is : <https://manifesto-project.wzb.eu>

Appendix Figure L1 depicts the number of active L and R politicians in our sample and their activity in each quarter. Left-wing politicians are more numerous than those on the right – there are 223 politicians classified as "Left", and 144 classified as "Right". We also classify leaders along the two dimensions simultaneously, to obtain four political types (LP, RP etc.). Appendix Table L1 reports the cross tabulation of the four types. Non-populist politicians in our sample tend to be left-wing, while populists are about equally split between left and right.

2.1.3 International Connections

The analysis that follows implicitly assumes that the distinctive rhetorical features of populism are similar across countries. To gauge the plausibility of this assumption, we study how connected are political leaders with others of the same type in different countries.

This can be ascertained by drawing on the Twitter follow check information (revealing who is followed by each leader), available for 306 politicians in our sample. Based on this information, we reconstruct follower networks among politicians. We then ask whether populist and non-populist leaders, or left and right leaders, are more likely to form ties (edges) with other political leaders of the same type across countries. Our key statistics here is a variant of the Coleman homophily index, which controls for the size of the political groups (there are less P and R than NP and L leaders in our sample). In particular, the index measures how much more likely it is that an individual connects with another individual of the same type (P vs NP, or L vs R) in a different country, compared to what would be expected by random chance. The index is given by:

$$H_i^{cross} = \frac{f_i - a_i}{1 - a_i}, \quad (1)$$

where f_i is the fraction of outgoing ties to nodes (Twitter accounts of leaders in our sample) who are of the same type as i but in a different country (relative to the total outgoing edges of i), a_i is the fraction of all nodes N that, for i , are of the same type but lie in a different country. When $H_i^{cross} > 0$ the node over-connects with same-type nodes in different countries, while $H_i^{cross} < 0$ means that it under-connects.

This index, displayed in Appendix Table E1, shows that populist lead-

ers have more international connections among themselves compared to non-populists. NP leaders heavily under-connect with same-type leaders abroad, while for P politicians the (unweighted) index is still negative but close to 0. Differences between P and NP leaders are statistically significant at 99% also when we weight nodes based on how much the politician is connected. Overall, P politicians connect internationally with other P politicians significantly more than their non-populist peers. This suggests populism is indeed akin to an international movement, more so than opposition to populism.

As documented in Table E2, populists are also significantly more likely to connect with other populists of the same L / R ideology, and among NP leaders, it is mainly RNP those who are internationally connected. Looking more broadly at overall international networks, we find that both LP and RP have around 1.3 to 1.5 more international connections, compared to non-populists, after controlling for network size. Within countries it is instead NP leaders, notably LNP, to be more connected (Appendix Table E4). The fact that individual leaders connect significantly more with leaders of the same type within countries than what would be expected by random chance is broadly supportive of our classifications.

2.2 Bigrams

All the tweets were translated into English using the open-access Python library googletrans, and then decomposed into bigrams. To do so, all special symbols were removed (punctuation marks, hyphens and apostrophes, emojis, user mentions and other hyperlinks, while hashtags were kept) and words were reduced to their stems.

The final dataset comprises around 40 mln bigrams, of which there are 15 mln unique bigrams. The vast majority of bigrams were used just once. Estimations drawing on such large number of bigrams would be computationally infeasible and in the remainder we restrict the analysis to the most frequent bigrams only. For our benchmark specification, only bigrams that were spoken in at least 10 calendar quarters and that have more than 24 occurrences in at least one quarter were used. We refer to this dataset restriction as to 10-25 (Figure J1 presents the histograms for these variables). This left us with about 6,000 unique bigrams and about 50% of the tweets. In section 6.3 we describe

the main results for different thresholds (5-50, 10-50).¹⁰

2.3 Topic classification

We classified bigrams using GPT and BERT into 19 topics chosen a priori (business, climate, energy, EU, foreign policy, public health, immigration, inequality, inflation, labor market, rights, security, taxes, trade, anti-establishment, personal communication¹¹, elections, self-promotion). We proceed in four steps:

1. we used the model GPT-5 mini to label a subset of tweets (100,000) into the 19 topics (see the Appendix for the prompt used in instructing ChatGPT);
2. we then trained BERT on these labeled tweets;
3. next we used the trained BERT to classify all 3.4 million tweets into topics with a threshold probability of .5;
4. finally, we classified bigrams on the basis of their relative recurrence in tweets belong to the different categories, calculating a simple measure of bigram specificity for each topic. In particular, bigrams exceeding the 40 % threshold in the ratio of the number of tweets containing bigram j on topic t to the number of tweets containing bigram j were assigned to the respective topics.

We additionally created a manually coded topic, *Specific Parties & Politicians*, which consists of bigrams that include the names of individual politicians or parties. In total, 4,483 bigrams (out of 6,144) were assigned to at least one topic (other than the residual one), and 3,705 were unique bigrams (a bigram can be assigned to more than one topic).

As a check of the validity of the classification provided by BERT we asked two master students to manually (and blindly) classify a sample of 510 tweets, which includes tweets on all of the topics. Appendix Table L4 displays the accuracy and F1 score of the different topics when compared to the manual classification.

¹⁰See the discussion on the construction of the dataset and the choice of these thresholds in Appendix J

¹¹Tweets only about personal communication such as greetings, weather or holidays

3 Estimation method

As explained by GST, the main challenge in comparing the language of different politicians is the much larger dimensionality of the available choice set, compared to the actually chosen words that we observe. In a finite sample, the use of different words by different politicians could be due to chance, not to a deliberate communication strategy. To cope with this problem, we rely on the methodology proposed by GST. Namely, we estimate the probability that a politician uses a specific bigram, correcting for finite sample bias.

3.1 Estimation

Suppose that there are only two political types, say populist ($P_i = 1$) or non-populist ($P_i = 0$) - below we allow for four political types. The utility of politician i from using bigram j in quarter t is:

$$U_{ijt}^{P_i} = \alpha_{jt} + X_{it}\gamma_j + \phi_{jt}P_i \quad (2)$$

where X_{it} are dummy variables for other features (not all time varying) of i , namely gender, higher education, being in government, being a candidate in the closest election, age. In addition, we also include fixed effects for country-by-quarter, and for quarterly distance from the closest election.¹² Finally, α_{jt} is a fixed effect for bigram j popularity in quarter t . The resulting probability that i uses bigram j in quarter t is:

$$q_{ijt}^{P_i} = \frac{\exp(U_{ijt}^{P_i})}{\sum_k \exp(U_{ikt}^{P_i})}, \quad P_i = 0, 1 \quad (3)$$

We observe the number of times that each bigram j is used by politician i in calendar quarter t , $\mathbf{c}_{it} = \{c_{ijt}\}$, and assume that this count variable has a multinomial distribution:

$$\mathbf{c}_{it} \sim MN(m_{it}, \mathbf{q}_{it}^{P_i}),$$

where $m_{it} = \sum_j c_{ijt}$ is the verbosity of i in quarter t and $\mathbf{q}_{it}^{P_i} = \{q_{ijt}^{P_i}\}$.

¹²Quarterly distance from the closest election captures election-specific bigrams, and the election window is defined as in the event studies that we estimate below. Lasso penalization implies that not all fixed effect coefficients are estimated, ensuring that collinearity with country-by-quarter fixed effects is not a concern.

For each bigram j we estimate parameters $\{\alpha_{jt}, \gamma_j, \phi_{jt}\}_{t=1\dots T}$ by minimizing the following penalized objective function, as in GST:

$$\sum_j \left[\sum_t \sum_i m_{it} \exp(\alpha_{jt} + X_{it}\gamma_j + \phi_{jt}P_i) - c_{ijt}(\alpha_{jt} + X_{it}\gamma_j + \phi_{jt}P_i) + \psi(|\alpha_{jt}| + ||\gamma_j||) + \lambda_j|\phi_{jt}| \right]$$

Two aspects of this procedure are worth mentioning. First, in order to make the model computationally feasible, we approximate the likelihood of the multinomial logit with the likelihood of a Poisson model $c_{ijt} \sim \text{Pois}(\exp(\mu_{it} + u_{ijt}))$ and the plug-in estimator $\hat{\mu}_{it} = \log(m_{it})$ for μ_{it} is used.

Second we make use of Lasso shrinkage for populist loadings $\lambda_j|\phi_{jt}|$. The value of λ_j is chosen separately for each bigram so as to minimize the corrected Akaike information criterion. Overall, this forces the parameter of interest ϕ_{jt} towards 0, reducing the finite sample error.¹³ For the other parameters, we set the penalty at $\psi = 10^{-5}$. This value controls the extent to which variation in bigram usage across individuals or over time is attributed to covariates rather than being absorbed by the populist dummies. Additionally, this penalty facilitates numerical convergence. We discuss the choice of the penalty in Appendix C.¹⁴

Using (2)-(3), we can estimate the probability that a politician of type (P_i, X_{it}) uses bigram j in quarter t , $\hat{q}_{ijt}^{P_i}$, for all j in our sample. The set of covariates is sufficiently rich that this estimated probability is specific to each individual politician.

Finally, note that implicit in (2)-(3) is the assumption that the propensity to use a specific bigram does not depend on which other bigrams were used by speaker i in session t , nor on the bigrams used by speakers of the opposite type, except through the bigram popularity fixed effects, α_{jt} , and the country

¹³In our baseline estimates, it turns out that ϕ_{jt} is non 0 for about 27% of bigrams per quarter.

¹⁴ Penalizing non-zero coefficients γ_j implies that around 22% of all the country by calendar quarter fixed effects are included on average. Given the importance of distinguishing between different countries and quarters, we thus also include in X_{it} separate fixed effects for each country and each calendar quarter. Without the penalty, this would be redundant, since these fixed effects would all be absorbed by the country by calendar quarter fixed effect. But our penalized estimation method implies that they can be separately estimated, although many of them are estimated to be zero, and the results of the event studies described below are not significantly affected by this inclusion (results available upon request).

by quarter fixed effects. These assumptions thus preclude a study of strategic interactions between competing politicians of opposite type.

3.2 Partisanship

Using $\hat{q}_{ijt}^{P_i}$, the posterior that an observer with neutral priors assigns to politician i being populist ($P_i = 1$), conditional on i having used bigram j in quarter t , is:

$$\hat{\rho}_{ijt} = \frac{\hat{q}_{ijt}^1}{\hat{q}_{ijt}^1 + \hat{q}_{ijt}^0} \quad (4)$$

Following GST, we measure the *partisanship* of politician i in quarter t as:

$$\hat{\pi}_{it} = \frac{1}{2} \sum_j [\hat{q}_{ijt}^1 \hat{\rho}_{ijt} + \hat{q}_{ijt}^0 (1 - \hat{\rho}_{ijt})] \quad (5)$$

This variable measures the average predictability of i 's type, given a single bigram. It is the average posterior probability with which an observer with neutral priors expects to correctly classify politician i as populist (P) or non-populist (NP), conditional on a single bigram. The average is computed over all possible bigrams, weighted by the estimated probabilities of their use by politician i in quarter t .

Note that any specific politician is either populist ($P_i = 1$) or non-populist ($P_i = 0$). The variable $\hat{\pi}_{it}$ averages the true type P_i and his "clone" $1 - P_i$ with neutral priors. In other words, partisanship measures the average predictability of two specific politicians with features (X_{it}) who are identical in everything but their political type. If P and NP always use the same bigrams, then $\hat{\pi}_{it} = 1/2$. If instead they always use different bigrams, then $\hat{\pi}_{it} = 1$.

As discussed in the previous section, we also measure partisanship between Right vs Left political types, rather than between P vs NP . This is done repeating the analysis described above, except that in (2)-(3) we replace the dummy variable P_i with a new dummy defined as $R_i = 1$ if i is right-wing, and $R_i = 0$ otherwise.

Moreover, we also consider a more granular partition of four political types, defined on two dimensions: populist / non-populist and right / left. To do so, when estimating q_{ijt} , we add to the RHS of (2) an additional linear term, $\varphi_{jt} R_i$, implying no interaction effects of the two political dimensions on word

usage.¹⁵ The appendix defines different measures of partisanship for this more granular model, corresponding to average predictability of the four types, or average predictability of the two types defined in each dimension (P/NP and R/L).

In our analysis, we also aggregate partisanship across politicians, or disaggregate it by and within topics. Specifically, average partisanship in calendar quarter t is

$$\bar{\pi}_t = \frac{1}{I_t} \sum_i \hat{\pi}_{it} \quad (6)$$

where I_t is the number of politicians active in quarter t . In Section 5 we measure time as distance d from national elections, and define $\hat{\pi}_d$ as average partisanship of active politicians who are at quarterly distance d from elections.

Finally, we define $\hat{\pi}_{it}^k$ as partisanship of i within topic k in quarter t as in (5), except that $\hat{q}_{ijt}^{P_i}$ is averaged only over bigrams belonging to topic k . Within-topic partisanship of i for all topics is the average of $\hat{\pi}_{it}^k$ over topics, weighting each topic k is by its frequency of use by i in quarter t . Partisanship of i between topics, instead, is constructed using $\hat{q}_{iJt}^{P_i} = \sum_{j \in J} \hat{q}_{ijt}^{P_i}$ computed with topics indexed by J (rather than bigrams) as units of speech. In all cases, $\hat{q}_{ijt}^{P_i}$ is always estimated over the entire sample of bigrams.

3.3 Chi-square measure of heterogeneity

A possible drawback of the partisanship indicators discussed above is that they measure heterogeneity of speech between two (ore more) predefined groups of politicians, lumping together other unobserved but potentially relevant political types. Thus, we also consider an alternative indicator of speech heterogeneity, based on the χ^2 statistics, that does not force us to take a stand on which are the relevant political types. This indicator is commonly used to quantify heterogeneity in other social contexts (see [Desmet, Ortuño-Ortín and Wacziarg \(2025\)](#)).

To cope with the finite sample problem highlighted by GST, we construct it using estimated probabilities of word usage, rather than observed word frequencies. Specifically, we estimate the probability that politician i with features X_{it} , uses bigram j in quarter t , $\hat{q}_{ijt}$, as described above, except that we

¹⁵To save on degrees of freedom, we impose the same Lasso penalty λ_j on the coefficients of both types, ϕ_{jt} and φ_{jt} .

impose $\phi_{jt} = 0$ in equation (2), thus neglecting political types.¹⁶

We measure how different is i 's rhetoric from that of the average politician in the same country by:

$$\hat{\chi}_{ict}^2 = \sum_j \frac{(\hat{q}_{ijt} - \hat{q}_{cjt})^2}{\hat{q}_{cjt}}, \quad (7)$$

where $\hat{q}_{cjt} = \frac{1}{n_c} \sum_{i \in c} \hat{q}_{ijt}$ is the probability of using bigram j averaged among all active politicians (n) in country c .

This indicator varies between 0 (if i behaves exactly like his country average) and a number greater than 1. Appendix B shows that $\hat{\chi}_{ict}^2$ can be interpreted as a measure of the average predictability of the politician having features (X_{it}), for an observer with neutral priors, given i 's word usage in quarter t . Note that, being based on individual deviations from the country average in the quarter, this indicator is sensitive to changes in the composition of active politicians across quarters.

As with partisanship, taking the average of $\hat{\chi}_{ict}^2$ over all active politicians in a given quarter results in a measure of aggregate rhetorical heterogeneity: $\hat{\chi}_t^2 = \frac{1}{n} \sum_i \hat{\chi}_{ict}^2$.

4 Partisanship over time

This section illustrates the main properties of our measures of polarization, and describes how they evolved over time and in reaction to key events common to most countries.

4.1 Most distinctive bigrams

We begin by examining which words are most indicative of populist versus non-populist rhetoric. Following the approach of GST, we assess the distinctiveness of each bigram based on its contribution to our partisanship measure.

¹⁶Recall that the vector X_{it} also includes some observed individual political features, such as being in or out of government and being a candidate in the closest election. Estimation of the relevant parameters continues to be penalized by the penalty $\psi = 10^{-5}$ and, as described in footnote 14, besides country by calendar quarter fixed effects we also include fixed effects for countries and calendar quarters. Using observed word frequencies, rather than the estimated probabilities $\hat{q}_{ijt}$, would result in a χ^2 statistics that is one order of magnitude larger than with the estimated $\hat{q}_{ijt}$.

Specifically, define the populist partisanship of a single bigram j used by politician i as the extent to which an observer with neutral priors would change her expected posterior that politician i is populist if that bigram was removed from the vocabulary, unbeknownst to the observer.¹⁷ The populist partisanship of bigram j in quarter t is the average of this measure across all politicians active in quarter t . Positive values of this measure denote bigrams typically used by populist politicians, negative values by non-populists.

Table 2 below lists the 10 most distinctive bigrams (i.e. with the largest absolute values) of populist and non-populist politicians in 2015 Q3 and 2020 Q1. These dates roughly correspond to two events that increased partisanship in our sample period, namely the inflow of refugees from Syria to Europe and the COVID shock.

Table 2: 10 most partisan populist and non-populist bigrams in 2015 Q3 and 2020 Q1

2015 Q3		2020 Q1	
Non-Populists	Populists	Non-Populists	Populists
Young People	Asylum Seeker	Prime Minister	Take Care
Prime Minister	Greek People	Young People	Good Morning
Climate Change	Good Morning	Public Health	Right Now
Good News	Press Conference	Good News	European Parliament
Sustainable Development	Via YouTube	Climate Change	Many People
Look Forward	Press Release	Will Continue	Will Never
The Anniversary	Can See	The Anniversary	Asylum Seeker
Work Together	Right Now	Work Together	Press Conference
Syrian Refugees	Illegal Immigration	Continue Work	Health Care
Will Work	Border Control	Member State	Bernie Sanders

Notes. The table reports the bigrams with the largest and smallest median values (taken over i) of the bigram partisanship measure ξ_{jit} for $t \in \{2015 \text{ Q3}, 2020 \text{ Q1}\}$ (footnote 17 defines ξ_{jit}). Larger median values correspond to populist bigrams, smaller median values – to non-populist bigrams.

Table 2 highlights the distinct language patterns employed by populist and non-populist politicians when discussing significant events, as well as in everyday communication. For example, during the refugee crisis (2015 Q3), non-populist politicians were more likely to refer to immigrants as *refugees*, often alongside references to other topics ("*climate change*", "*sustainable de-*

¹⁷ Following GST, the populist partisanship of bigram j used by politician i in quarter t is:

$$\xi_{jit} = \frac{1}{2} - \frac{1}{2} \sum_{k \neq j} \left(\frac{q_{ikt}^1}{1 - q_{ijt}^1} + \frac{q_{ikt}^0}{1 - q_{ijt}^0} \right) \rho_{ikt}$$

velopment"). Populist politicians instead frequently emphasized phrases such as *"asylum seekers"*, explicitly mentioned that immigration is *illegal*, or mentioned *"border control"*. During the COVID-19 epidemic (2020 Q1), the bigram *"public health"* was distinctive of non-populist discourse, while populist politicians emphasized connection with people with the bigram *"take care"*, and tended to use language less specific to the pandemic, and continued to focus on immigration.

Everyday communication is also revealing of politicians' types. Non-populist politicians tend to prioritize more substantive policy agenda and policymaking process, while populist politicians focus on direct communication and establishing connections with the electorate. In both periods non-populists tend to talk about *"climate change"*, *"young people"*, and to refer to the bodies of establishment (e.g. *"prime minister"*). Instead, populists' everyday bigrams are more direct and less policy oriented – e.g. *"good morning"*, *"via youtube"*.

4.2 Average partisanship

We now turn to our core measures of partisanship, which summarize the distinctiveness of each politician by aggregating the numerous bigrams they use into a single metric. Figure 2 displays the estimated measures of polarization, averaged over individuals per calendar quarter.¹⁸

Several points are worth highlighting. First, two indicators – χ^2 (dark gray line) and populist / non-populist partisanship (pink line, $\bar{\pi}_{t,PNP}$) – move closely together and both peak during the European immigration crisis (2014–2016) – their correlation coefficient is 0.78. Notably, this pattern emerges even though the χ^2 measure does not rely on any predefined group assignments, and it refers to within-country heterogeneity. This co-movement suggests that overall polarization in the data mirrors the dynamics captured by populist partisanship. We interpret this as suggestive evidence that the populist dimension plays a central role in political competition within our sample. Note however that, at the level of individual politicians, populist partisanship, $\hat{\pi}_{it}$, and χ^2_{ict} are much less tightly correlated (their correlation coefficients is only 0.14). Hence, these two indicators are different measures of rhetorical

¹⁸Appendix Figure L2 reports the results with estimated confidence intervals for average partisanship. Since standard nonparametric bootstrap is invalid for lasso regression, we use subsampling instead, though it is considerably less efficient when the number of units (politicians) is relatively small.

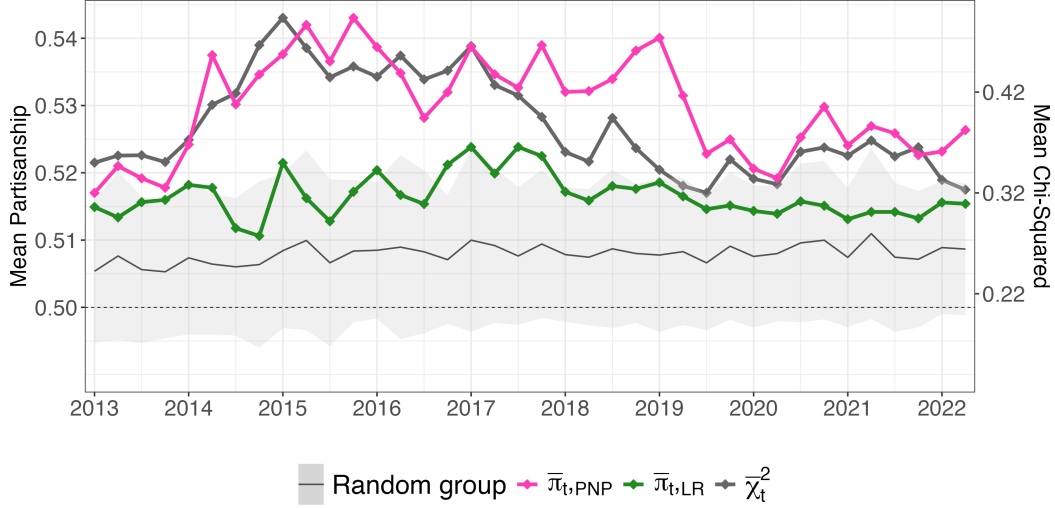


Figure 2: Average partisanship and χ^2 by calendar quarter

Notes. $\bar{\chi}_t$ refers to average χ_{ict}^2 per calendar quarter t in the model with no political types; $\bar{\pi}_{t,LR}$ refers to average partisanship defined in the model with two political types: left and right; $\bar{\pi}_{t,PNP}$ refers to average partisanship defined in the model with two political types: populist and non-populist. Random group refers to average partisanship with two randomly defined political types; the shaded area is the 95% confidence interval. Partisanship is defined as in Equation 5. Chi-Squared measure is defined as in Equation 7. Average values are defined as in Equation 6.

polarization.

The prominence of populist polarization over left-right polarization is confirmed by the fact that average left-right partisanship (green line, $\bar{\pi}_{t,LR}$) remains consistently lower than for the populist dimension across the entire period. This could suggest that the left-right divide is less central to political competition in many countries, or indicate that politicians are less aligned internationally on this dimension, and they lack a common rhetoric or agenda across borders. This could also be due to significant variation in how left-right positions are defined and understood across national contexts, providing further evidence of the dominant importance of the populist divide in our international sample of politicians.

The level of populist partisanship estimated in our sample is notably higher and more volatile than in GST’s analysis of U.S. Congressional speeches over the past 35 years.¹⁹ As described below, our estimates of $\hat{\pi}_t$ imply that, upon observing 5 bigrams (the typical content of a single tweet), the posterior of a politician being or not populist reaches about 0.6 during the refugee cri-

¹⁹In GST’s data, partisanship based on their preferred penalized estimator—which is comparable to the one we use—peaks at approximately 0.512 in 2010.

sis. By contrast, the posterior of a speaker being Republican or Democrat in GST’s sample, conditional on 5 bigrams, peaks at 0.55. The larger and more volatile partisanship in our sample is likely to reflect multiple factors. Tweets represent a high-frequency, low-cost, and highly flexible medium of communication, allowing politicians to react immediately to unfolding events and tailor messages to specific audiences. In contrast, Congressional speeches are more static, tend to follow more standardized formats, with topics and rhetoric constrained by institutional norms. Additionally, our estimates are computed at the quarterly level, and politicians differ in their levels of activity across time. These factors jointly contribute to the elevated and more volatile patterns of populist partisanship observed in our setting.

Finally, observed partisanship measures are substantially higher than those estimated when we randomly divide politicians in two groups reflecting the observed frequencies of true P /NP types in our sample, repeating this procedure 100 times (the black line is the average, the shaded area is the 95% confidence interval). This random assignment serves as a benchmark to quantify the extent of remaining finite-sample bias in our measure of average partisanship. By construction, partisanship within random groups should be around 0.5. In our case, average random partisanship values are generally just below 0.51, and significantly above 0.5.²⁰ This is probably due to the relatively small number of politicians in our sample, and to the possible presence of unobserved political types. We discuss this point more extensively in the next section and in Appendix K.

We now investigate these time fluctuations more thoroughly, focusing on the effect of the immigration crisis and the COVID-19 pandemic on the measures of populist partisanship, within the topics of immigration and public health respectively.

4.3 Topic partisanship during crises

4.3.1 The 2015 refugee crisis

In 2015 about 1.3 million Syrian refugees escaped to Europe. As shown in Figure 2, populist partisanship peaks in 2015; moreover, as shown in Table 2 many of the most partisan bigrams in 2015 Q3 refer to immigration. In order

²⁰For each calendar quarter, we reject the null hypothesis that mean partisanship equals 0.5 at the 5% significance level.

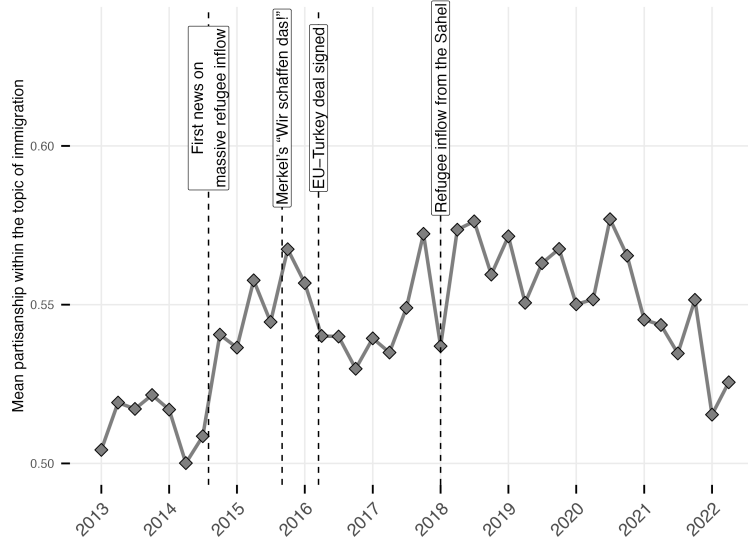


Figure 3: Partisanship within the topics of immigration

Notes. The figure reports the average P–NP partisanship by calendar quarter, computed using the sample of bigrams that belong to the immigration topic and restricting the sample to European countries. Partisanship is defined as in Equation 5. Average values are defined as in Equation 6.

to assess the impact of this shock on political rhetoric, we focus on European countries, and on the immigration topic.

Figure 3 displays the average populist partisanship of speech within the topic of immigration for the sample of European countries. To calculate this measure we use the same definition of partisanship as in Equation 5, but restrict the analysis only to the subsample of immigration bigrams. Hence, we are estimating the predictability of the politician’s type based on the used bigrams within the topic of immigration.

Figure 3 reveals an increase in partisanship within the topic of immigration beginning in the fourth quarter of 2014, coinciding with the onset of the mass inflow of Syrian refugees. This upward trend continues, peaking in late 2015 – around the time of Chancellor Merkel’s decision to open Germany’s borders to Syrian refugees. Partisanship remains high throughout the remainder of the sample period, even after the initial stages of the immigration crisis; the magnitude of partisanship within the immigration topic significantly exceeds that observed for aggregate partisanship (see Figure 2). Together, these patterns support prior research identifying immigration as a highly salient and mobilizing topic, especially for populist politicians.

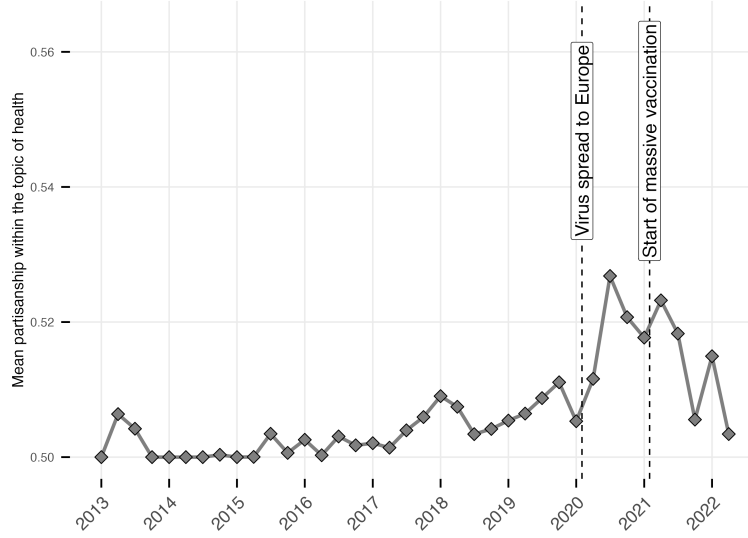


Figure 4: Partisanship within the topics of public health

Notes. The figure reports the average P–NP partisanship by calendar quarter, computed using the sample of bigrams that belong to the public health topic. Partisanship is defined as in Equation 5. Average values are defined as in Equation 6.

4.3.2 The COVID shock

The COVID pandemic is another shock common to several countries, that polarized public opinion and politicians over the severity of lockdowns and the vaccine requirements. We thus repeat the same analysis of the previous subsection, but focusing on all the countries in the sample and on bigrams within the topic of public health. The results are presented in Figure 4.

The results indicate that public health was not a particularly divisive topic prior to the onset of the COVID-19 pandemic. With the emergence of the pandemic, however, populist partisanship within this topic increased sharply, before declining again as the crisis subsided. Nonetheless, the absolute levels of partisanship remain substantially lower than those observed for the topic of immigration. This suggests that, while debates around lockdown measures and vaccination became politically salient in some contexts, they did not map as clearly onto the populist versus non-populist divide.

5 Electoral competition

Different political types may use different language because they have different intrinsic ideologies, or because they pursue different strategies to win

voters’ support. To disentangle these different motives, we now study how our measures of heterogeneity evolve during the electoral cycle. If different language mostly reflects an opportunistic strategy to win votes, then heterogeneity should rise as the election date approaches and decline afterwards. If instead differences in language are mainly driven by the politician’s intrinsic motivation, then if anything they should be larger after the election.

In this section we document an increase in rhetorical heterogeneity before the elections dates, quantify the effect, and propose mechanisms that could explain our findings. We exploit the fact that election dates differ across countries, as shown in Appendix Table L3. The election date refers to national parliamentary elections for parliamentary regimes. For presidential and mixed systems, the election date refers to all parliamentary elections and presidential elections in which more than two candidates are present in our sample of politicians (France, Mexico, the US, Slovenia, Portugal, Finland, Czech Republic).²¹

5.1 Average measures of heterogeneity over the electoral cycle

Figure 5 plots the average measures of rhetorical polarization by quarterly distance d from the election, $d=0$ being the election quarter. On average, populist vs non-populist partisanship increases in the run-up to the election, reaching especially high values two quarters prior and remaining high through the election quarter, and dropping very fast right after. In other words, as the election date approaches, the language of populist and non-populist politicians becomes more distinctive. Appendix Figure L3 shows that, when conditioning on calendar-quarter fixed effects, the pre-electoral increase in partisanship is even more pronounced. The shaded pink area represents 90% point-wise confidence intervals of populist vs non-populist partisanship based on sub-sampling. Appendix D discusses how the confidence intervals are constructed.

Appendix Figure L6 plots average partisanship values for the Left-Right division over the electoral cycle. Left-right partisanship exhibits almost no electoral dynamics. This confirms the dominance of the populist vs non-populist dimension of electoral competition in our sample of countries, and we mostly

²¹We exclude elections at levels different than national. In our politician–election panel, the mean of the indicator for participating in an election (either directly as a candidate, or indirectly as the member of a party participating in the election) is 0.6.

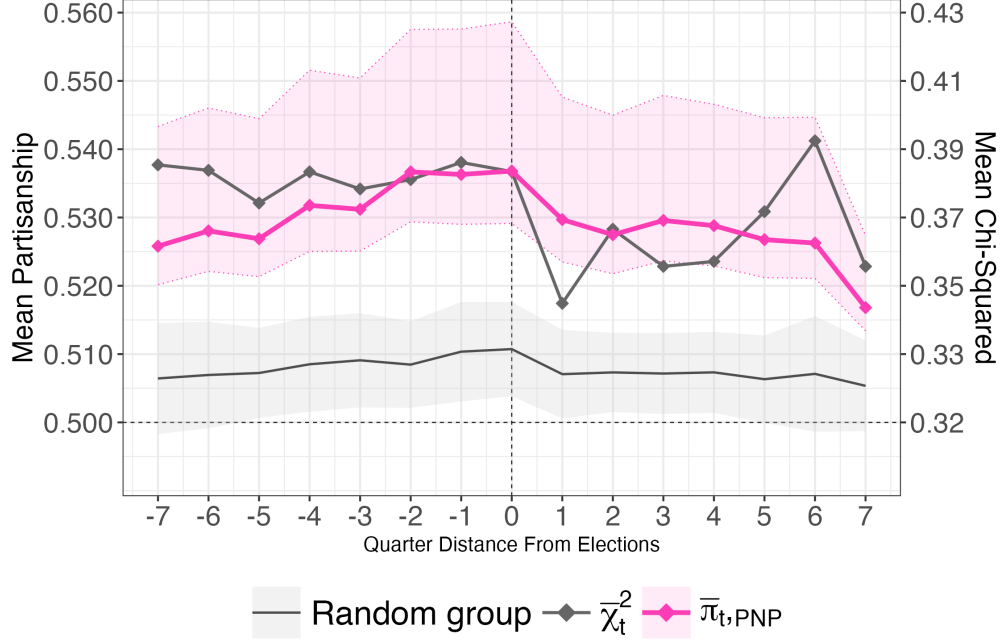


Figure 5: Average estimated χ_{it}^2 and π_{it} per quarter difference from elections.

Notes. $\bar{\chi}_t$ refers to average χ_{it}^2 per quarter distances t in the model with no political types; $\bar{\pi}_{t,PNP}$ refers to average partisanship defined in the model with two political types: populist and non-populist. Shaded light pink area behind $\bar{\pi}_{t,PNP}$ is 90% pointwise CI based on subsampling. Random group refers to average partisanship with two randomly defined political types; the shaded area behind the random group is the 95% confidence interval. Partisanship is defined as in Equation 5. Chi-Squared measure is defined as in Equation 7. Average values are defined as in Equation 6.

focus on this dimension in the following parts of the paper.

The χ^2 measure also rises as the election approaches and declines thereafter, with a similar pattern to populist partisanship, indicating a general rise in rhetorical heterogeneity within countries, irrespective of political affiliation.

Note that all these patterns refer to aggregate measures of polarization, so they are also affected by changes in the composition of active politicians near the elections. We deal with this issue in Subsection 5.3 below, where we estimate an event study at the level of individual politicians.

The solid line at the bottom of the figure depicts partisanship for randomly grouped politicians, and the shaded gray area is the 95% confidence interval. The estimates for the populist vs non-populist partisanship are well above the random group benchmark. Importantly, the difference between populist partisanship and the random group benchmark increases before the election and peaks at the election quarter, suggesting that the populist dimension of rhetorical polarization has a significant electoral component.

As already noted with regard to Figure 2, the random group benchmark is significantly above 0.5, and here it also exhibits some pre-electoral dynamics. To explain these patterns, Appendix K runs Monte Carlo simulations using synthetic samples of politicians and bigram usage. We explore two issues: we vary the number of politicians in the sample, and we allow for the presence of unobserved political types. Both aspects are relevant. As the number of politicians increases, the random group benchmark converges towards 0.5. But the presence of electorally relevant unobserved groups also matters and can explain the pre-electoral patterns. The reason is that, in the presence of prominent unobserved types, their bigram usage may be spuriously attributed to random groups by the penalized estimator, resulting in partisanship estimates that systematically deviate from the 0.5 benchmark (although we find this effect to be small in practice). Appendix K provides further details on the simulation design and results.²²

5.2 Informativeness of speech over the electoral cycle

Recall that partisanship measures the average predictability of a politician’s type, conditional upon observing a single bigram. According to Figure 5, on average for populist vs non-populist politicians this predictability increases from 0.525 five quarters before the election to 0.535 in the election quarter, and then declines by about the same amount two quarters after the election. How can this magnitude be interpreted?

To address this question, we extend the analysis to multiple bigrams and compute a measure of *speech informativeness* conditional on observing n bigrams, for different values of n . Following the logic of the model in Section 3.1, and as in GST, for each politician i and calendar quarter t we randomly draw N bigrams from a multinomial distribution with bigram choice probabilities given by the estimated frequencies, $\hat{\mathbf{q}}_{it}^{P_i} = \{\hat{q}_{ijt}^{P_i}\}$ (P_i denotes a politician i ’s true type). Let $\hat{q}_{int}^{P_i}$ denote the frequency of use of the n -th drawn bigram for politi-

²²This pattern is also likely to reflect the properties of the penalized estimator, which is bounded below by 0.5, but could exceed 0.5 in the random sample of politicians due to its non-linearity and convexity. The reason is that, despite the Lasso penalization, in a relatively small sample some bigrams will, by pure chance, correlate with the random political label. We verify that this is indeed the case by implementing the leave-out estimator proposed by GST, which is an out-of-sample estimator, and as such it is much less subject to the positive bias due to Jensen’s inequality. The leave-out estimator recovers the 0.5 benchmark for random-group partisanship even when unobserved types are present, in line with the findings of GST.

cian i in calendar quarter t . Then, for any sequence length $n \in \{1, 2, \dots, N\}$, we compute the following measure of speech informativeness conditional on the observed sequence:

$$\eta_{itn} = \frac{\hat{q}_{int}^{P_i} \eta_{itn-1}}{\hat{q}_{int}^{P_i} \eta_{itn-1} + \hat{q}_{int}^{1-P_i} (1 - \eta_{itn-1})}, \quad (8)$$

where $\eta_{it0} = \frac{1}{2}$ for all i and calendar quarters t (i.e. as before, an observer starts with a neutral prior).²³ Averaging η_{itn} over politicians i in a given quarter t we obtain a measure of speech informativeness η_{tn} . To obtain a similar measure at a given quarter distance d from the election, we take all the politicians-quarter observations at a given quarter distance d , and average $\eta_{id(i,t)n}$ over i to obtain η_{dn} .

Figure 6 presents the average results of repeating this procedure 10 times through a Monte Carlo simulation *for each politician* in each session. The vertical axis measures speech informativeness about the politician being populist or non-populist, the horizontal axis measures the number of bigrams. One tweet corresponds to about 5 bigrams (the average length of a bigram is about 20 characters and the average length of a tweet is about 100 characters, although tweets also include many rarely used bigrams that we have excluded from the analysis). Different colors correspond to different distances from the election date.

Upon reading a single tweet two years before the election, an observer with neutral priors would raise its posterior about the politician’s populist type by about 7 percentage points (from 0.5 to about 0.57). During the election quarter, the increase in predictability from reading one tweet jumps to almost 12 percentage points (from 0.5 to 0.62), almost doubling the informativeness of a single tweet. After reading two tweets (10 bigrams) during the election quarter, predictability increases by 17 percentage points (from 0.5 to 0.67), in contrast to an increase of 12 percentage points (from 0.5 to 0.62) if two tweets are read two years before the election, an increase of about 50%. Note that the curves are concave, implying that the largest marginal increases in informativeness are obtained after just a few bigrams.

We can compare these estimates to the informativeness of text at two specific dates: 2015 Q3 (the dark dotted line), during the first acute phase of

²³This measure corresponds to the expected posterior probability of correctly identifying politician i ’s type after observing n bigrams, starting with a neutral prior.

an immigration crisis; and 2020 Q1 (the light dotted line), the beginning of the COVID pandemic. These two dates refer to one of the lowest and the highest values of populist partisanship in our sample respectively. The texts from 2015 Q3 (immigration crisis) are slightly less informative about populist partisanship than the texts written during the electoral quarter. This is as expected, as immigration is one of the most divisive issues between populists and non-populists.²⁴ This is not true for the COVID-19 shock, and the informativeness of text in this period more closely follows the informativeness 2 years before the elections.

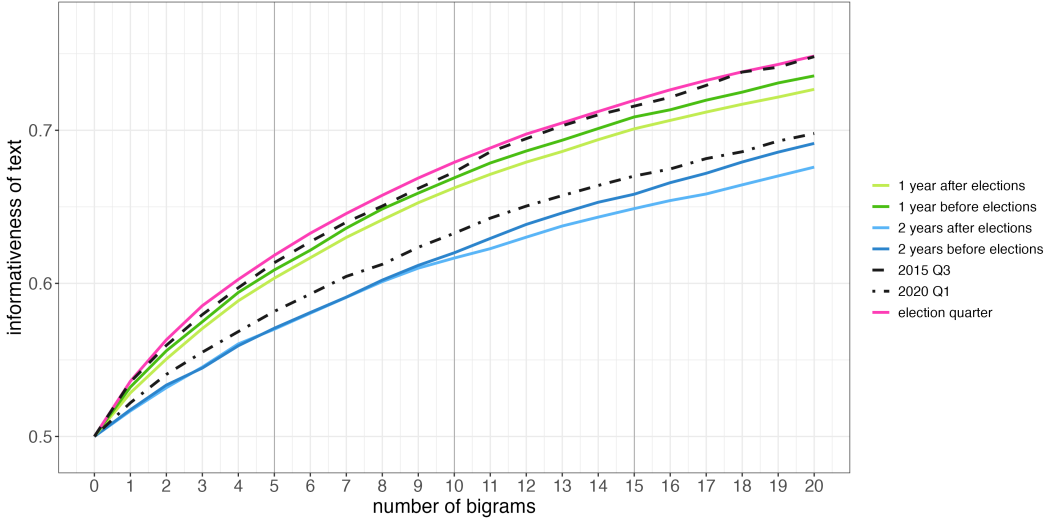


Figure 6: Expected posterior belief of an observer with a neutral prior after reading a given number of bigrams

Notes. Figure plots the expected posterior probability that an observer with a neutral prior assigns to politician i 's true populist affiliation after hearing n bigrams (η_{itn}), averaged over i . η_{itn} is defined as in Equation 8. Estimated bigram-usage probabilities $\hat{\mathbf{q}}_{it}$ from our benchmark specification are used. Vertical lines indicate approximately one, two, and three tweets (five, ten, and fifteen bigrams, respectively).

5.3 Event study

The estimates displayed in Figure 5 could reflect the coincidence of elections with other relevant events, or changes in the composition of active politicians. To control for these potential confounders, we exploit staggered election dates across countries and estimate an event study with politician (γ_i), calendar date (δ_t) and country by election-window (El_{it}) fixed effects, namely:

²⁴The rise in populist partisanship during the electoral cycle is of similar magnitude to that observed in the 2015 refugee crisis, with average values around 0.54.

$$Y_{it} = \sum_{d(i,t)} \beta_{d(i,t)} D_{d(i,t)} + \gamma_i + \delta_t + El_{it} + \epsilon_{it}, \quad (9)$$

The dependent variable is a measure of rhetorical polarization for politician i (either π_{it} or χ_{it}^2), and $D_{d(i,t)}$ are indicators for quarterly distance from the election - which in our baseline specification include 3 quarters before and after the election as well as the election quarter. The election-window is defined as the country-quarters closest to each election. Election-window fixed effects, El_{it} , are included to capture country specific shocks occurring during the election-window - results are robust to dropping El_{it} . Thus, the estimated coefficients $\beta_{d(i,t)}$ measure the average effect on rhetorical polarization of being at distance $d(i,t)$ from the election, relative to the average quarters at a greater distance from the election within each election-window, for the same politician. Errors are clustered at the (closest) election level - the treatment variable of interest.²⁵

Figure 7 shows the event studies for the χ_{it}^2 measure of polarization (left panel) and the populist partisanship measure π_{it} (right panel) - note that the scales of the vertical axis are different. The event studies confirm that the language used by politicians becomes more polarized two quarters before the election and remains so until the election date. Once the election is over, however, rhetorical heterogeneity declines to its average level outside of the election window. In terms of magnitudes, the estimated rise in populist partisanship during the election quarter relative to non-election quarters is about 10% of one standard deviation of the populist partisanship measure across the full sample (0.0065/0.058). The corresponding estimate in the χ_{it}^2 model is instead about 3% of one standard deviation in the full sample (0.013/0.51).

In Appendix Table L2, we report a Wald test to assess whether the estimated coefficients β for the post-electoral periods are statistically different from the estimate for the quarter immediately preceding the election ($d(i,t) = -1$), and the electoral quarter ($d(i,t) = 0$). The results indicate that, for partisanship, the null hypothesis of equality is rejected at the 5% significance level for all post-election quarters except $d(i,t) = 3$, for both

²⁵That is, clustering is among observations sharing the same fixed effect El_{it} . Thus, we allow arbitrary correlations in residuals among observations closest to the same country-election date. To define El_{it} , for each country and calendar quarter, we find which national election is closest to the middle day of that quarter. This creates an electoral-cycle fixed effect by country, with no ties because “closest” is defined using the exact middle day of the quarter.

quarters $d(i, t) = -1$ and $d(i, t) = 0$. In contrast, for the χ_{it}^2 model, the hypothesis of equality cannot be rejected between pre-electoral and all post-electoral quarters. The estimates for the electoral quarter are different from those in $d(i, t) = 1$ at the 1% significance level. Thus, the electoral pattern of rhetorical polarization is larger and statistically more significant for populist partisanship, π_{it} , than for the politically agnostic measure χ_{it}^2 .

Appendix Figure L4 depicts the event study for left-right partisanship. We do not find a sizable effect of rhetorical polarization before the elections for these two groups.

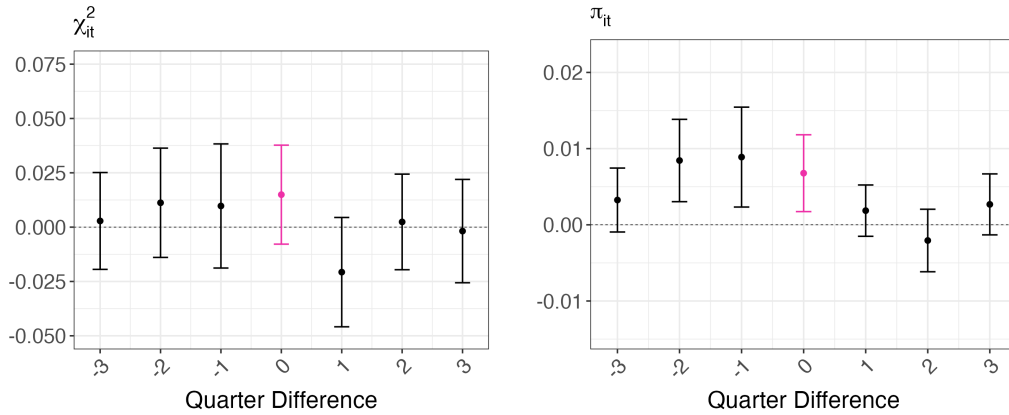


Figure 7: Event studies for main measures of segregation

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from specification 9. Outcome variables are chi-squared measure χ_{it}^2 (left panel), and populist partisanship π_{it} (right panel). Errors are clustered at the elections level, 95% CI are reported. Average χ_{it}^2 in the sample is 0.376, average π_{it} is 0.523. Partisanship is defined as in Equation 5. Chi-Squared measure is defined as in Equation 7.

Appendix Figure L5 depicts the event study on partisanship when the probabilities q are estimated allowing for four political types, as in Appendix A. The left-most panel refers to populist vs non-populist partisanship, the center panel to left-right partisanship, and the right-most panel refers to partisanship between four types (note that the scale here is smaller, since with four types partisanship ranges between 0.25 and 1). The pre-election pattern is present in all panels, but it is weaker for the left-right division.

Overall, these findings suggest that, while there is a general increase in rhetorical polarization before the elections, this increase is more pronounced for the populist dimension of political competition. From here on we focus in more detail on this dimension only.

6 Mechanisms and Robustness

What mechanisms can explain this pre-electoral divergence in the language of populist vs non-populist leaders? Changes in the composition of the audience are unlikely to play a significant role. Leaders’ communication on Twitter/X is typically addressed to external audiences, not to party insiders, and the number of politicians’ followers does not fluctuate around election dates (Van Kessel et al., 2020).

On the other hand, as the election approaches, gaining votes is likely to become a more relevant goal, relative to other motives such as expressing personal convictions, or reacting to external events.²⁶ Particularly during an electoral campaign, tweets may be used strategically by politicians to: i) communicate where they stand on a specific policy issue; ii) draw voters’ attention to specific topics; iii) persuade voters by changing their beliefs in specific ways. Although we cannot unpack these different motives, in this section we explore more specific aspects of partisanship during the electoral cycle.

6.1 Different electoral systems

In standard models of electoral competition where candidates announce their policy positions ahead of the elections, like Downs’ or probabilistic voting, candidates’ incentives to converge depend on the number of competing candidates. Two vote maximizing candidates who compete for the same voters have strong incentives to converge. With more than two candidates or with the threat of entry by a third candidate, however, convergence is not assured (e.g. Palfrey (1984)). Although we have emphasized the distinction between only two political types, several countries in our sample have more than two parties. Could this account for the observed divergence?

To address this question, we split the sample in two groups of elections: (i) all national elections held in countries with plurality rule and in presidential regimes plus all presidential elections, where competition is typically between

²⁶Severin-Nielsen et al. (2025) interviewed Danish MPs about how their communication strategies differ between election and routine periods. They highlight that, while in routine periods communication is often developed on an ad hoc basis, during an electoral campaign it is much more thoroughly and coherently thought out, also because more resources are devoted to it. This is how a Danish MP described his communication during a routine period: “It is a bit more random than you would think. (. . .) It’s very focused on issues of the day (. . .). And there are a lot of posts that are based on an impulse”.

two serious parties or candidates; (ii) parliamentary elections in proportional systems, that typically have several competing parties and candidates in the same districts.²⁷

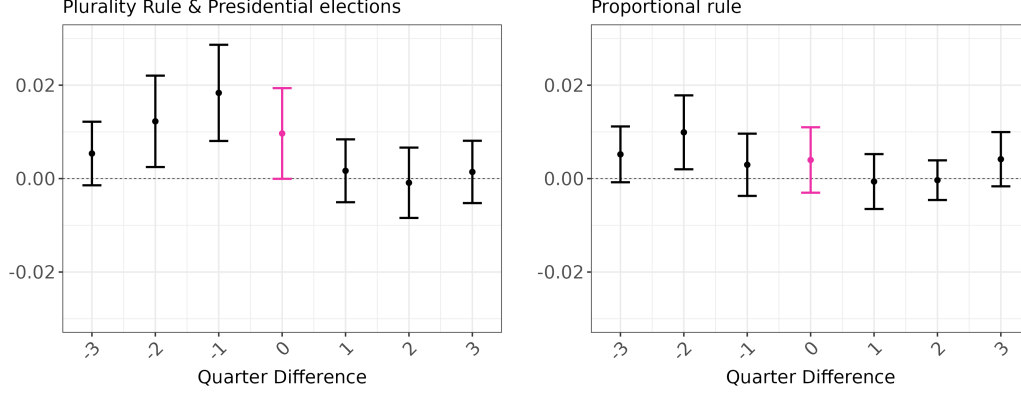


Figure 8: Partisanship event study estimates for elections under different electoral systems.

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from specification 9 separately for the two subsamples of elections held under different institutional rules. Left Panel plots the results for all national elections held in countries with plurality rule and in presidential regimes, as well as for all presidential elections, where competition is typically between two serious parties or candidates; Right Panel – for parliamentary elections in proportional systems, that typically have several competing parties and candidates in the same districts. Outcome variable is partisanship, π_{it} , as defined in Equation 5. Errors are clustered at the elections level, 95% CI are reported.

Figure 8 reproduces the event study separately for these two groups of elections. While the pre-electoral polarizing effect is present in both groups of elections, it is much larger for plurality rule and presidential elections. Hence, it is unlikely to be driven by competition between more than two parties. To the extent that tweets are used to communicate policy positions, we have to look for other opportunistic motives that induce divergence in two-party systems. One explanation could be that voters are more attentive to the communication of their favored leader, so that each party targets a different group of voters rather than the same swing voters as his opponent. In this case, electoral incentives induce divergence in policy platforms, as leaders find it optimal to energize their core supporters (Glaeser, Ponzetto and Shapiro

²⁷Group (i) includes all elections held in Australia, the UK, the US, France, Brazil, Mexico and Hungary, plus presidential elections with more than one active candidate in our sample of political leaders in Finland, Poland and Slovenia. In Brazil presidential and parliamentary elections are always held in the same dates. In Mexico three congressional elections in our sample were held without a contemporaneous presidential election; we nevertheless include all Mexican elections in group (i) because they tend to be dominated by electoral competition for president. Hungary has a mixed system that strongly favors the largest party.

(2005)) or to please the voters who are more attentive to their communication (Gennaioli and Tabellini (2025)). Note that asymmetric voters’ attention can also create incentives to diverge if leaders’ communication aims to change voters’ beliefs, rather than to inform voters of their policy positions - cf. Gennaioli and Tabellini (2025).

6.2 Partisanship between and within topics

As discussed above, tweets are likely to be used to draw attention to specific issues, or to persuade voters through issue framing or other rhetorical arguments, and not just to communicate policy positions. To explore these other aspects of communication strategies, we classify bigrams into topics and calculate partisanship values *between* and *within* topics. We discuss the construction of these measures in the Section 3.2. All the estimations in this section are done using the results from the benchmark specification.

6.2.1 Decomposition of average partisanship

Figure 9 plots the within- and between-topics partisanship measures over the electoral cycle. Both within- and between-topics partisanship rise as the election date approaches and decline after the election.

Within-topic partisanship is consistently higher than between-topic partisanship. This implies that it is, on average, easier to infer a politician’s type from a specific bigram within a topic than from the topic label itself. This pattern is consistent with Table 2, which reports the most distinctive bigrams by politician type: for example, both populists and non-populists discuss immigration, but the particular bigrams they use on this topic are different and predictive of a politician’s type.

Figure 10 illustrates the event study for partisanship between and within topics. Both increase before the election and peak at the election date, confirming the patterns described in Subsections 5.1 and 5.2. The pre-electoral rise in within-topic partisanship means that, as the election approaches, populist vs non-populist politicians are more likely to use different bigrams when referring to the same issues. This could reflect policy positioning or persuasion. The rise in between topic partisanship before the election is consistent with the idea that populist and non-populist politicians seek to draw attention to different sets of issues.

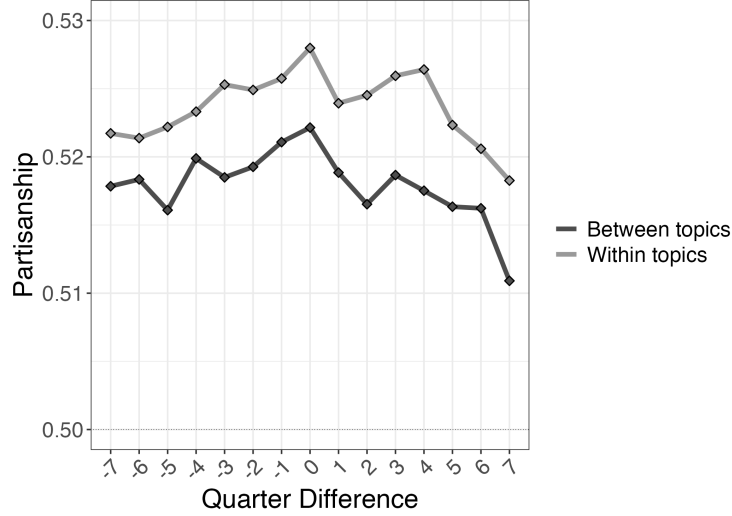


Figure 9: Measures of segregation over the electoral cycle

Notes. Between-topic partisanship uses topic as a unit of speech; within-topic measure considers only bigrams within the topics as units of speech, and weights topics by their frequencies of use. These measures were calculated based on predicted frequencies of use of specific bigrams from the benchmark model fit. Subsection 3.2 defines these measures.

Both types of partisanship complement each other, so that an aggregate rise in partisanship around the elections is higher than both between and within estimates. Note that about 40% of bigrams are not included in the analysis depicted in Figure 10, since they do not belong to any classified topic. These non-classified bigrams too exhibit some pre-electoral cycle in partisanship, as described in Appendix L19.

6.2.2 Policy or Non-Policy?

What might explain the increasing partisanship between and within topics as elections approach? One of populism’s defining traits is its *thin-centered ideology*, which the literature connects to the absence of clearly defined policy positions among populist politicians, compared with non-policy elements in their rhetoric (Mudde and Kaltwasser, 2012). This implies that, among the topics we consider, some are likely to be typical of populist politicians, while others are more characteristic of non-populists. Moreover, one can expect different electoral dynamics in the frequency of tweeting about these topics over the electoral cycle. For example, non-populist politicians may campaign more on policy issues they perceive as advantageous, or seek to counter populist narratives, while populist politicians may further downplay policy dimensions

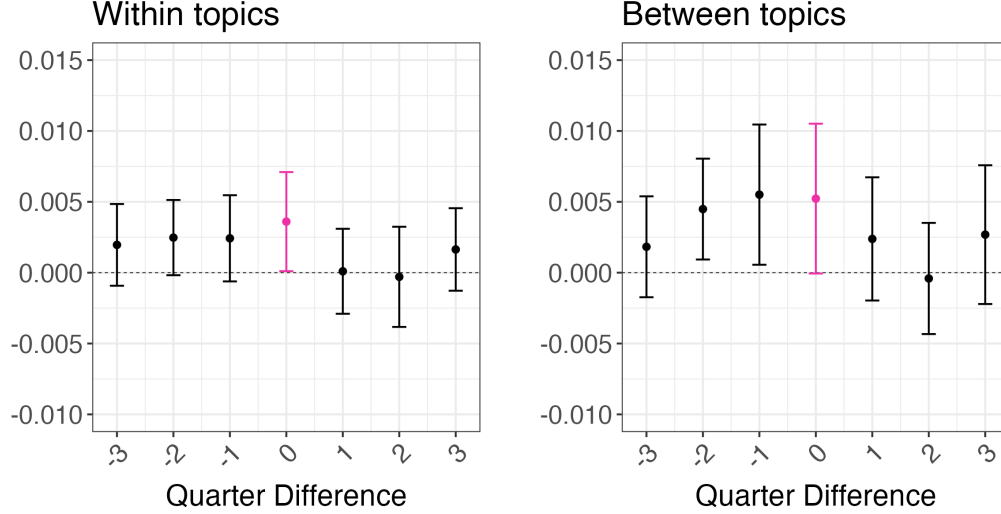


Figure 10: Event studies for the between- and within-topics partisanship

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from specification 9, using within-topic partisanship (left panel) and between-topic partisanship (right panel) as outcome variables. Errors are clustered at the elections level, 95% CI are reported. Subsection 3.2 defines within- and between-topic partisanship.

as elections approach, shifting attention even more toward non-policy issues.

We first ask: Is there a clear distinction between populist and non-populist topics, with the latter being more policy-oriented? To address this, we rely on the variable ρ_{iJt} defined in equation (4), namely the posterior that speaker i is populist, conditional on having spoken about topic J . A topic is considered populist if the average of ρ_{iJt} over the entire sample exceeds 0.5, and non-populist otherwise. Average values of ρ_{iJt} by topic are reported in Table 3.²⁸

Apart from immigration and post-COVID health, the topics most typical of populist politicians are non-policy oriented (mentions of specific parties and politicians, self-promotion, anti-establishment content, and elections). By contrast, the topics most typical among non-populist politicians are predominantly policy-oriented, with climate, foreign policy, and inflation among the most distinctive. Taken together, the rhetorical patterns we observe in the data are consistent with existing findings in the literature.

We next turn to study the evolution of populist vs non-populist topics over the electoral cycle, separately for all populist topics lumped together (i.e. all topics J such that $\rho_{iJ} > 0.5$ in Table 3.) and all non-populist topics (i.e. all topics J such that $\rho_{iJ} < 0.5$). We estimate an event study specification similar to Equation (9), except that now the dependent variable is ρ_{iKt} , namely the

²⁸Looking at median ρ_{iJt} , rather than average, gives very similar results.

Table 3: Average populist distinctiveness by topic

Topic	Average ρ_{iJt}
Immigration	0.5589
Specific parties & politicians	0.548
Self-promotion	0.527
Anti-establishment	0.513
Public Health (post-COVID)	0.506
Elections	0.505
Inequality	0.493
Labour market	0.487
Personal Communication	0.485
Security	0.485
Taxation	0.483
European Union	0.481
Energy policy	0.481
Rights	0.478
Business & industry	0.478
Public Health (pre-COVID)	0.477
Trade policy	0.477
Inflation	0.477
Foreign policy	0.475
Climate & Environment	0.468

Notes. Average topic distinctiveness $\bar{\rho}_J$ is the mean of ρ_{iJt} over individuals i and quarters t for topic J . For each individual i and quarter t , we first compute frequencies of using bigrams on topic J , q_{iJt} . Distinctiveness is then defined as

$$\rho_{iJt} = \frac{q_{iJt}^{\text{Pop}}}{q_{iJt}^{\text{Pop}} + q_{iJt}^{\text{Non-Pop}}},$$

where q_{iJt}^{Pop} and $q_{iJt}^{\text{Non-Pop}}$ are the predicted frequencies for the same covariates under populist and non-populist affiliation, respectively.

posterior that i is populist, conditional on speaking about the set of topics K , where K denotes the set of either populist or non-populist topics listed in Table 3.

Figure 11 presents the results. The green solid line in both graphs represents the mean value of ρ_{iKt} across the entire sample, for populist and non-populist topics. The estimates should be interpreted in relation to this line. The green shaded area represents the range between 0.5 (i.e., when a topic is addressed with equal frequency by populists and non-populists) and the mean of ρ_{iKt} . Values within this shaded area suggest that the language of both types of politicians is converging, approaching a value of 0.5. In contrast, values outside this area indicate that the language is diverging, moving away from 0.5.

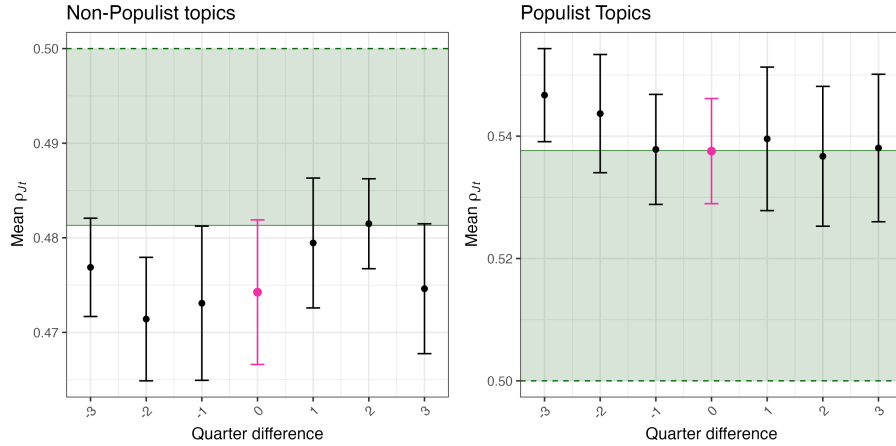


Figure 11: Predicted frequencies of speech over electoral cycle by sets of topics.

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from specification 9 with ρ_{iJt} as an outcome variable. Errors are clustered at the elections level, 95% CI are reported. The estimates are re-centered relative to the mean level ρ_{iJt} and should be interpreted relative to this line. Measure ρ_{iJt} takes values above 0.5 for topics more frequently used by populist politicians, and below 0.5 for topics more frequently used by non-populist politicians; and the conclusions about divergence/convergence should be drawn accordingly. The green shaded area on the graphs implies an area of *convergence* – estimates in this area mean topic is becoming *less telling*. Estimates out of this area imply *divergence*.

The results suggest that, closer to elections, non-populist topics become even more typical of non-populist (left panel), while there is no pronounced electoral pattern in populist topics. In terms of magnitude, the estimated effect in pre-electoral quarters is about 10% of one standard deviation of the ρ_{ijt} for the non-populist topics across the full sample. The dynamics for non-populist topics are consistent with several interpretations – for example, non-populist politicians may campaign more intensively on policy issues as elections ap-

proach, or they may abandon these topics, but less so than populist politicians do. The flat pattern for populist topics is consistent with a scenario in which both groups of politicians increasingly adopt populist rhetoric. The analysis of the drivers of speech distinctiveness in the next section provides further support for this interpretation.

Overall, the results of partisanship between and within topics supports our hypothesis that populist and non-populist politicians focus their followers' attention on fundamentally different sets of topics, operating along the extensive margin of rhetoric, that is, the choice of the topics.

6.3 Robustness

We now discuss a number of possible concerns with our estimates of partisanship, and show their robustness.

Snap elections About 20% of elections in our data are snap elections - i.e. called prematurely before the end of the legislature. Although snap elections do not happen without premonitions, they may have a shorter pre-electoral cycle in communications. Moreover, unobserved confounding events could cause both a rise in rhetorical polarization and the occurrence of snap elections. To allow for this, we estimate the specification separately for elections held on time and the elections held ahead of time. Appendix Figure L7 shows that indeed snap elections induce a shorter and less precisely estimated cycle of polarization (also because of the smaller sample), which is more focused on the electoral quarter itself. Nevertheless, results remain unaffected for the regular election dates.

Being a candidate Although all politicians in our sample are prominent political leaders with strong stakes in the election outcomes, not all of them stand as candidates in all elections. On the one hand, one would expect that the leaders who are also candidates have sharper electoral incentives, and so their rhetoric should exhibit a stronger pre-electoral cycle. On the other hand, rhetorical polarization could be the result of politicians diverging in their immediate goals – some of them are seeking reelections, while others are continuing with their usual agenda, or supporting running candidates. To explore these conjectures, we analyze how the results differ among politicians who are candidates in given elections, and those who are not. Appendix Figure

L8 shows that the pre-election cycle is starker among the running candidates, in line with the idea that it reflects their stronger electoral incentives, rather than different goals of politicians in more heterogeneous situations.

English speaking countries Given that our analysis is in multinational context and we translate all the tweets to English, one might suspect that our results are driven more by English-speaking countries. This could also potentially explain why the electoral cycles are more pronounced in the group of plurality rule and presidential elections, where the fraction of English speaking countries is higher. But Appendix Figure L9 shows that this is not the case: partisanship increases in both groups of countries, with the effect actually being more pronounced in non-English-speaking countries.

Time window We also perform robustness checks, with regard to the length of the time window of the event study. In Appendix Figure L10 we consider longer windows of 4, 5, 6 quarters (rather than 3) around the elections. In principle this could matter because it changes the non-treated observations. But the main results are confirmed, and the election cycle generally begins two quarters before the election, as in the baseline estimates.

Heterogeneous treatment effects Could the results of the event studies be contaminated by heterogeneous effects driven by staggered adoption? Appendix Figure L11 shows that our results are robust to this concern, as they are almost identical when implementing the staggered difference-in-difference estimations suggested by Sun and Abraham (Sun and Abraham, 2021).

Verbosity Politicians are more verbose closer to the elections (Appendix Figure L12). Could this electoral pattern in verbosity, rather than the language used, explain our divergence results? We don't think so. First, the Poisson model that we estimate incorporates verbosity as a parameter. So, unless there is a direct link between verbosity and speech polarization, this should not be a concern (Gentzkow, Shapiro and Taddy, 2019). Second, the model is estimated at the level of calendar quarters, and the political-type coefficient ϕ_{jt} is not individual specific by design. Hence, verbosity of individual politicians cannot directly translate into higher partisanship estimate.

Self-selection into being active A potentially more relevant issue is the extensive (rather than intensive) margin of politicians’ activity. Quarters closer to the election could have a larger number of active politicians. By construction, our measure of segregation is based on estimated non-zero coefficients ϕ_{jt} from the Lasso-like optimization problem. Hence, all else equal, the larger is the number of speakers active during the quarter, the more there are non-zero estimated bigram-specific loadings on the populist dummy variable, and the larger is the estimated value of partisanship.

This issue too is not a concern in our sample, however, for several reasons. First, we estimate partisanship over calendar quarters. Hence, the number of speakers could be causing a problem only if those calendar quarters with larger number of speakers are also more likely to be election quarters. This is hardly the case (see figure L13 in the Appendix).

Second, while our results indicate the profound effect on partisanship in the electoral quarter and in the two quarters immediately preceding it, the number of active speakers does not differ much within a six months window around the elections time (Appendix Figure L14). The reason is that our sample consists almost exclusively of high-ranking politicians, most of whom use Twitter on a regular basis to communicate with the electorate. This is illustrated in Appendix figure L15, which provides the distribution of the number of politicians corresponding to each maximum absolute quarterly distance from the elections. There are almost no politicians who tweet exclusively around electoral quarters.

Third, in the event study we explicitly control for all quarters in the election window, thus capturing mean partisanship levels during these quarters. Hence, indirectly, this captures the popularity of the session too.

Fourth, to make sure that our results are not driven by different numbers of speakers per calendar quarters, we perform a robustness check, keeping 220 speakers²⁹ in each calendar quarter at random. Note that most sessions in our sample have more than 220 active speakers, as shown in Figure L13, so this simulation could give less precise and smaller estimates of partisanship. The event study results of such robustness analysis are in Appendix Figure L16 and confirm the main results

²⁹220 is the minimum observed number of politicians (2013 Q1) that allows us to perform selection without replacement

Sample of bigrams We also applied different thresholds for bigrams selection. Appendix Figure L17 shows that the event studies are unaffected.

Populist classification We demonstrate that our findings remain robust under an alternative classification of populist politicians, in which ambiguous cases are coded as non-populists rather than as populists, as in the main specification – Appendix Figure L18.

7 Electoral cycles in populist distinctiveness

Populists are often new challengers, who try to fill representation gaps by catering to dissatisfied voters (Guenther, 2024). Non-populists, instead, are more likely to be mainstream politicians who appeal to moderate voters (Table 1). If so, one could expect that populists try to distinguish themselves through novel communication strategies and by raising new issues, while non-populists may be forced to imitate the challengers to limit their electoral gains (Vries and Hobolt, 2020). Our measure of partisanship cannot tell us which politician type is mostly responsible for electoral divergence. Nevertheless, we can ask whether, as the election approaches, language is becoming more or less distinctively populist. Of course, this could be due to populists becoming more extreme in their rhetoric, to imitation by the non-populist, or both.

To address this question, we focus on the frequency with which each politician uses distinctively populist words, namely:

$$y_{i,t} = \sum_{j \in J} \hat{q}_{ijt}^{P_i} \xi_j \quad (10)$$

where $\hat{q}_{ijt}^{P_i}$ are the probabilities of bigram usage estimated in earlier sections, and ξ_j is a measure of populist distinctiveness of bigram j , which only varies across bigrams, and not across politicians or quarters. ξ_j is obtained from the variable $\xi_{i,j,t}$ already defined in footnote 17, as discussed below. Recall that $\xi_{i,j,t}$ has both direction and magnitude: positive values of $\xi_{i,j,t}$ indicate populist bigrams, negative values indicate non-populist ones. Its magnitude reflects the strength with which the bigram is associated with populist or non-populist speech. Thus, the variable $y_{i,t}$ measures the average populist distinctiveness of the bigrams used by politician i in quarter t . This indicator varies over time and across politicians only due to frequency of usage of distinctively populist

bigrams.

To construct the time invariant measure of bigram populist distinctiveness, ξ_j , we proceed as follows. First, when estimating the probability that leader i uses bigram j in quarter t , that goes into the computation of $\xi_{i,j,t}$, we constrain parameter ϕ_{jt} (capturing the influence of political type on bigram choice) to be constant over time. All other time varying parameters that enter into the estimation of $\hat{q}_{ijt}^{P_i}$ are included in the same form as in the baseline estimation discussed above. For each bigram j , this gives us a measure of populist distinctiveness, ξ_{ijt} , where the effect of political type on bigram usage is constrained to be the same in all quarters.

This measure of bigram populist distinctiveness is still very noisy, however, because it varies across politicians and quarters due to all other determinants of bigram usage, besides political type³⁰. To compute a time invariant and population wide measure of populist distinctiveness that only varies at the bigram level, ξ_j , for each bigram j we compute the median value of ξ_{ijt} over its entire distribution across quarters and politicians. We use the median rather than the average to avoid the influence of extreme quarters or politicians with very distinctive speech, but results are similar when using instead the average of ξ_{ijt} over i and t . This procedure gives us a measure of how distinctively populist is each bigram, aggregating populist distinctiveness over time and across politicians. From an observer’s perspective, this measure is more natural than ξ_{ijt} , as the public typically perceives certain phrases as more or less populist without reference to specific individuals or time periods.³¹

We then estimate an event-study specification identical to Equation (9), but with y_{it} as the outcome variable. We start by defining the set J in the summation above as the full set of 6144 bigrams. The results are presented in Figure 12. The estimates indicate that, as the election approaches, political language becomes increasingly populist. The effect is sizable: in the election quarter, speech distinctiveness rises by about 0.17 standard deviations relative to quarters outside of the elections window, with a somewhat smaller effect

³⁰The variance decomposition exercise shows that a substantial share of the variation in this measure of distinctiveness is explained by individual fixed effects. Our estimates indicate that individual fixed effects account for 33% of the total variation, while calendar-quarter fixed effects explain about 4%.

³¹Note that ξ_{ijt} depends on the speech and characteristics of politicians from both groups. Aggregating it over time and across politicians allows us to construct a unique measure, that can be directly mapped to the usage of specific bigrams across sessions, making our analysis easier to interpret.

in the pre-electoral quarter. By contrast, point estimates for other quarters are close to zero and statistically indistinguishable from zero at conventional confidence levels.

Because the distinctiveness measure ξ_j is time-invariant, observed trends capture shifts in the relative usage of bigrams over the electoral cycle. The results therefore suggest that politicians tend to use more populist bigrams closer to elections, substituting away from non-populist ones.

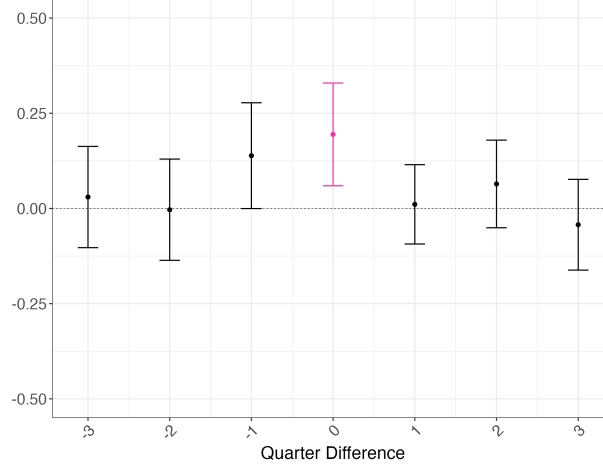


Figure 12: Event study estimates for the distinctiveness of the speech

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from specification 9. The outcome variable y_{it} is standardized and defined as in Equation 10. Errors are clustered at the elections level, 95% CI are reported.

Upon examining the results in Figure 12, two natural questions arise: (i) which sets of bigrams drive the effect, and (ii) which types of politicians (populist or non-populist) are responsible for this trend? To address the first question, we focus on the most extreme bigrams in both tails of the ξ_j distribution. Specifically, we replicate the analysis defining J as the set of only the most populist bigrams (those with ξ_j above the 80th percentile of the ξ_j distribution) and the most non-populist bigrams (those with ξ_j below the 20th percentile). Changes in distinctiveness for bigrams from the upper tail capture shifts in the use of populist language closer to elections, while changes for bigrams from the lower tail reflect shifts in the usage of non-populist language.

The results of this analysis are shown in Figure 13. They indicate no electoral pattern in bigram usage for the most populist bigrams (left panel). In contrast, there is a clear electoral pattern for the most non-populist bigrams: speech distinctiveness in the lower tail is becoming increasingly more populist before and during the election quarter. This suggests that politicians on

average use non-populist bigrams less frequently. The effect of elections is substantial, with lower-tail speech being approximately 0.20 standard deviations more populist in the electoral quarter relative to the quarters outside of the studied window (± 3 quarters).

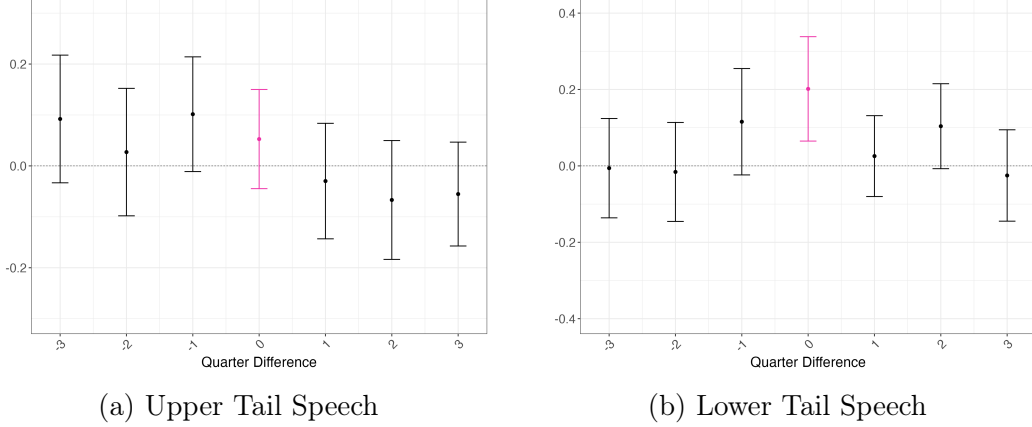


Figure 13: Event study estimates for the distinctiveness of tails of speeches
Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from specification 9 on two subsamples: the left panel uses the most populist bigrams (with ξ_j above the 80th percentile of the ξ_j distribution), and the right panel uses the most non-populist bigrams (with ξ_j below the 20th percentile). The outcome variable y_{it} is standardized and defined as in Equation 10. Errors are clustered at the elections level, 95% CI are reported.

In Appendix Figures L20, L21 we present the results of the same analysis for populist and non-populist politicians separately. The results for both groups of politicians are similar and mirror the main results presented here.

A plausible interpretation of these findings is that non-populist politicians seek to emulate the rhetorical style of populists in the lead-up to elections. That is, closer to election day non-populist politicians are more likely to adopt a direct and emotionally resonant style of communication—features that are characteristic of populist rhetoric (see Section 4.1, which highlights the linguistic features most strongly associated with populist and non-populist discourse).³² Note that this result parallels our findings in Section 6.2.2, where we observe no clear electoral dynamics in the distinctiveness of populist topics. If non-populist politicians increasingly emulate populist rhetoric, the distinctiveness of populist topics will remain relatively stable over time.

These findings are not in contrast with the divergence results described above. In this part of the analysis we assume that phrase distinctiveness ξ_j

³²These distinctively non-populist bigrams often refer to specific institutions, other countries, or terms used by political elites, such as *prime minister*, *head of state*, *government must*, *launching new*, and *member state*.

is fixed across politicians and time. This implies that each phrase is fixed to be either populist ($\xi_j > 0$) or non-populist ($\xi_j < 0$) across the whole sample. However, in the main analysis each phrase’s distinctiveness is being referenced relatively to the phrases used by the politician i in a quarter t (see footnote 17 for a definition of this measure). For instance, the phrase *close border*, which is on average populist, could be making some populist politician seem more non-populist, if all the rest of their phrases are more populist (e.g. such phrases as *asylum seeker*, *border control* among others, see Table 2). Hence, what the results in this section highlight, is that both groups of politicians tend to substitute on average more non-populist phrases with on average more populist phrases.

Moreover, the fact that language is becoming more populist just before the election does not imply that politicians from opposite groups start using more similar bigrams. In general, there is no clear relationship between distinctiveness and partisanship. Both groups’ speech could become more populist, but the bigrams they use remain different. Given the staggered nature of elections and the use of calendar-quarter estimates, it is likely that similar developments in the distinctiveness measure across groups over the electoral cycle would still result in a higher measure of partisanship.

8 Concluding Remarks

This paper studies how electoral competition shapes the communication strategies of political leaders. Using high-frequency data from Twitter/X for 367 leaders in 21 democracies over 2013–2022, we document a large electoral cycle in political rhetoric. As elections approach, political language of populist vs non-populist politicians becomes more polarized, with rhetorical polarization reverting to the average after the election. This pattern is stronger in plurality rule and presidential elections, suggesting that it is not due to a large number of parties, and it is observed irrespective of the country language.

This electoral cycle suggests that rhetorical divergence is not driven solely by fixed ideological differences. Rather, it reflects an opportunistic response to electoral incentives, as politicians strategically differentiate their communication when electoral stakes are highest. The fact that differentiation involves both how the same issues are framed as well as which topics are emphasized suggests that this is not only about valence or party names, but it also concerns

policy views and political ideology.

The populist dimension plays a central role in this process. Compared to the conventional left–right divide, populist versus non-populist polarization is both more pronounced and more responsive to the electoral cycle. As elections approach, non-populist topics—largely policy-oriented except for immigration and health policy after COVID—become even more diagnostic of non-populist politicians, while within-topic framing differences between populists and non-populists also widen. At the same time, all politicians, regardless of type, adopt a more populist style of communication near elections, reducing the use of abstract and institutional language characteristic of non-populist rhetoric.

These patterns are consistent with a setting in which populist vs non-populist politicians do not compete for the same swing voters, but instead target and seek to mobilize different groups of voters. This is in line with the results of [Allcott et al. \(2025\)](#), who show that in the 2020 presidential election most ads in a large sample of Facebook and Instagram users were targeted towards parties’ own supporters. If political communication is not uniformly attended to, but reaches different audiences, then competing candidates have an incentive to diverge in their policy platforms and to polarize their core supporters with their political propaganda ([Glaeser, Ponzetto and Shapiro, 2005](#) and [Gennaioli and Tabellini, 2025](#)).

These findings point to three important directions for future research.

First, is pre-electoral differentiation a general feature of political communication, or is it only present on Twitter/X? The convergence results by [Di Tella et al. \(2025\)](#) on party and candidates official positions suggest that social media may play a special role. But this question deserves further attention, both with regard to other media, and because the convergence result in [Di Tella et al. \(2025\)](#) reflects a comparison of final elections with US primaries and French first round elections, whereas we compare pre-electoral and electoral quarters with communication after the election or in more distant quarters.

Second, and related to the above, greater attention should be paid to how platform design and targeting technologies affect political competition. Much of the literature on social media has emphasized how echo-chambers or the diffusion of specific emotional content may fuel political extremism, with the evidence pointing in both directions ([Melnikov, 2024](#) and [Allcott et al., 2025](#)). Our findings highlight a second potentially important mechanism, that operates on the supply side. By lowering the cost of selective communication,

social media may fundamentally alter candidates' incentives, encouraging differentiation and polarizing propaganda.

Third, although populism is generally characterized by policy positions endorsing nationalism and intolerance of multiculturalism (Guriev and Papaioannou, 2022), our analysis suggests that populism is also associated with more cross-country interactions between political leaders than traditional politics. Such international spillovers in narratives and communication over identity politics are largely unexplored and deserve to be investigated more thoroughly.

References

- Alesina, Alberto. 1988. "Credibility and policy convergence in a two-party system with rational voters." *The American Economic Review* 78(4):796–805.
- Allcott, Hunt, Matthew Gentzkow, Levy Ro'ee et al. 2025. The Effects of Political Advertising on Facebook and Instagram before the 2020 US Election. Technical report National Bureau of Economic Research.
- Ash, Elliott, Massimo Morelli and Richard Van Weelden. 2017. "Elections and Divisiveness: Theory and Evidence." *The Journal of Politics* 79(4):1268–1285.
- Aslanidis, Paris. 2018. "Measuring populist discourse with semantic text analysis: an application on grassroots populist mobilization." *Quality & Quantity* 52(3):1241–1263.
- Battiston, Giacomo. 2022. "Rescue on Stage: Border Enforcement and Public Attention in the Mediterranean Sea." *Available at* <https://www.dropbox.com/sh/ct8qe4x71l7hjry/AABxUJxsYmQy1Tea2> .
- Besley, Timothy and Stephen Coate. 1997. "An economic model of representative democracy." *The Quarterly Journal of Economics* 112(1):85–114.
- Bonica, Adam, Kasey Rhee and Nicolas Studen. 2025. "The Electoral Conse-

- quences of Ideological Persuasion: Evidence from a Within-Precinct Analysis of US Elections.” *Available at SSRN* .
- Bonomi, Giampaolo. 2024. “Divide and diverge: Polarization incentives.” *Available at <https://ssrn.com/abstract=4990250>* .
- Brady, David W., Hahrie Han and Jeremy C. Pope. 2007. “Primary Elections and Candidate Ideology: Out of Step with the Primary Electorate?” *Legislative Studies Quarterly* 32(1):79–105.
- Canes-Wrone, Brandice, David W Brady and John F Cogan. 2002. “Out of step, out of office: Electoral accountability and House members’ voting.” *American Political Science Review* 96(1):127–140.
- Cassell, Kaitlen J. 2023. “The Populist Communication Strategy in Comparative Perspective.” *The International Journal of Press/Politics* 28(4):725–746.
- Chatterjee, Arindam and Soumendra Nath Lahiri. 2011. “Bootstrapping lasso estimators.” *Journal of the American Statistical Association* 106(494):608–625.
- Desmet, Klaus, Ignacio Ortuno-Ortín and Romain Wacziarg. 2025. On the Measurement of Social Heterogeneity. Technical report National Bureau of Economic Research.
- Di Tella, Rafael, Randy Kotti, Caroline Le Pennec and Vincent Pons. 2025. “Keep your enemies closer: strategic platform adjustments during US and french elections.” *American Economic Review* 115(8):2488–2528.
- Dommett, Katharine, Samuel A Mensah, Junyan Zhu, Tom Stafford and Nikolaos Aletras. 2025. “Is there a permanent campaign for online political advertising? Investigating partisan and non-party campaign activity in the UK between 2018–2021.” *Journal of Political Marketing* 24(2):143–161.
- Fiva, Jon, Oda Nedregård and Henning Øien. forthcoming. “Group Identity

- and Parliamentary Debates.” *Journal of Politics* .
- Fowler, Anthony and Shu Fu. Forthcoming. “Do primary elections exacerbate congressional polarization?” *Journal of Politics* .
- Funke, Manuel, Moritz Schularick and Christoph Trebesch. 2023. “Populist leaders and the economy.” *American Economic Review* 113(12):3249–3288.
- Gennaioli, Nicola. and Guido Tabellini. 2025. “Identity Politics.” *Econometrica* 93(6):1937–67.
- Gentzkow, Matthew, Jesse M Shapiro and Matt Taddy. 2019. “Measuring group differences in high-dimensional choices: method and application to congressional speech.” *Econometrica* 87(4):1307–1340.
- Gibson, Rachel K. 2020. *When the nerds go marching in: How digital technology moved from the margins to the mainstream of political campaigns*. Oxford University Press.
- Glaeser, Edward L., Giacomo Ponzetto and Jesse Shapiro. 2005. “Strategic Extremism: Why Republicans and Democrats Divide on Religious Values.” *The Quarterly Journal of Economics* 120(4):1283–1330.
- Guenther, Laurenz. 2024. “Political Representation Gaps and Populism.” *Available at SSRN 4230288* .
- Guriev, Sergei and Elias Papaioannou. 2022. “The Political Economy of Populism.” *of Economic Literature* 60(3):753–832.
- Guriev, Sergei, Nikita Melnikov and Ekaterina Zhuravskaya. 2021. “3g internet and confidence in government.” *The Quarterly Journal of Economics* 136(4):2533–2613.
- Hall, Andrew B. 2015. “What happens when extremists win primaries?” *American Political Science Review* 109(1):18–42.

- Hall, Andrew B and James M Snyder Jr. 2015. "Candidate ideology and electoral success." *Manuscript, Stanford University* .
- Hill, Seth J. 2017. "Changing votes or changing voters? How candidates and election context swing voters and mobilize the base." *Electoral Studies* 48:131–148.
- Kamada, Yuichiro and Fuhito Kojima. 2014. "Voter Preferences, Polarization, and Electoral Policies." *American Economic Journal: Microeconomics* 6(4):203–36.
- Klein, Ezra. 2020. *Why We're Polarized*. Simon & Schuster, Inc.
- Manacorda, Marco, Guido Tabellini and Andrea Tesei. 2022. "Mobile internet and the rise of political tribalism in Europe." *Available at https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4411693*.
- Meguid, Bonnie M. 2005. "Competition Between Unequals: The Role of Mainstream Party Strategy in Niche Party Success." *American Political Science Review* 99(3):347–359.
- Melnikov, Nikita. 2024. "Mobile internet and political polarization." *Available at SSRN 3937760* .
- Mudde, Cas and Cristóbal Rovira Kaltwasser. 2012. *Populism in Europe and the Americas: Threat or corrective for democracy?* Cambridge University Press.
- Norris, Pippa. 2019. "Varieties of populist parties." *Philosophy & Social Criticism* 45(9-10):981–1012.
- Norris, Pippa. 2020. "Measuring populism worldwide." *Party Politics* 26(6):697–717.
- Palfrey, Thomas R. 1984. "Spatial Equilibrium with Entry." *The Review of Economic Studies* 51(1):139–156.

- Persson, Torsten and Guido Tabellini. 2002. *Political economics: explaining economic policy*. MIT press.
- Polborn, Mattias K. and Jr. Snyder, James M. 2017. “Party Polarization in Legislatures with Office-Motivated Candidates.” *The Quarterly Journal of Economics* 132(3):1509–1550.
- Prummer, Anja. 2020. “Micro-Targeting and Polarization.” *Journal of Public Economics* 188:1–21.
- Rooduijn, Matthijs. 2019. “State of the field: How to study populism and adjacent topics? A plea for both more and less focus.” *European Journal of Political Research* 58(1):362–372.
- Severin-Nielsen, Majbritt Kappelgaard, Sanne Kruikemeier, Jakob Ohme and Kristian Gade Kjellmann. 2025. “From permanent campaign to permanent communication: parliamentarians’ cross-media practices in routine and election times.” *Journal of Information Technology Politics* 14:1–17.
- Sun, Liyang and Sarah Abraham. 2021. “Estimating dynamic treatment effects in event studies with heterogeneous treatment effects.” *Journal of econometrics* 225(2):175–199.
- Van Kessel, Patricia, Regina Widjaya, Sono Shah, Aaron Smith and Adam Hughes. 2020. “The congressional social media landscape.” Available at <https://www.pewresearch.org/internet/2020/07/16/1-the-congressional-social-media-landscape/>.
- Vries, Catherine E de and Sara Hobolt. 2020. *Political entrepreneurs: The rise of challenger parties in Europe*. Princeton University Press.
- Zhang, Lehan, Gabriele Gratton, Pauline Grosjean and Hasin Yousaf. 2025. “Adding Fuel to the (Gun) Fire: How Politicians Polarize the Public Debate.”.

Appendix

A Partisanship with four political types

Suppose there are four political types, defined on two dimensions: populist / non-populist and left / right. Hence, rewrite (2) as:

$$U_{ijt}^{P_i R_i} = \alpha_{jt} + X_{it} \gamma_j + \phi_{jt} P_i + \varphi_{jt} R_i$$

with $P_i = 0, 1$, $R_i = 0, 1$. We thus have four political types, Populist - Right, Non-Populist - Right, and so on, and estimate a vector of probabilities for each type, $q_{it}^{P_i R_i} = \{q_{ijt}^{P_i R_i}\}$.

The posterior probability that i is of type, say, Populist - Right, with neutral priors and conditional on bigram j , is:

$$\hat{\rho}_{ijt}^{11} = \frac{\hat{q}_{ijt}^{11}}{\hat{q}_{ijt}^{11} + \hat{q}_{ijt}^{10} + \hat{q}_{ijt}^{01} + \hat{q}_{ijt}^{00}}$$

and similarly for the other types.

Partisanship of i in this more granular model is the weighted average of these posterior probabilities across all bigrams for all possible types, weighted by the probability of their occurrence, namely:

$$\hat{\pi}_{it4} = \frac{1}{4} \sum_j \sum_{P=0,1; R=0,1} \hat{q}_{ijt}^{PR} \hat{\rho}_{ijt}^{PR}$$

This measures the predictability of i 's (granular) type by averaging his true type $P_i R_i$ with all his possible "clones" with neutral priors. If all types always use the same bigrams, then they are undistinguishable and $\hat{\pi}_{i4} = 1/4$. If instead they always use different bigrams, then $\hat{\pi}_{i4} = 1$.

We can also estimate the predictability of P vs NP , when there are 4 political types as:

$$\hat{\pi}_{it4}^{P/NP} = \frac{1}{4} \sum_j [(\hat{q}_{ijt}^{11} + \hat{q}_{ijt}^{10}) \hat{\rho}_{ijt}^{1\circ} + (\hat{q}_{ijt}^{01} + \hat{q}_{ijt}^{00})(1 - \hat{\rho}_{ijt}^{1\circ})]$$

where $\hat{\rho}_{ijt}^{1\circ} = \hat{\rho}_{ijt}^{11} + \hat{\rho}_{ijt}^{10}$ is the posterior that i is populist ($P_i = 1$). This expression for partisanship measures the average predictability of i 's type being P or NP , when there are four possible political types. Similarly, we can define

right-left partisanship when there are four political types as:

$$\hat{\pi}_{itA}^{R/L} = \frac{1}{4} \sum_j [(\hat{q}_{ijt}^{11} + \hat{q}_{ijt}^{01})\hat{\rho}_{ijt}^{\circ 1} + (\hat{q}_{ijt}^{10} + \hat{q}_{ijt}^{00})(1 - \hat{\rho}_{ijt}^{\circ 1})]$$

where now $\hat{\rho}_{ijt}^{\circ 1} = \hat{\rho}_{ijt}^{11} + \hat{\rho}_{ijt}^{01}$.

B Links between χ^2 and partisanship

By Bayes rule, the posterior that the speaker has features X_{it} , rather than being some other politician of the same country c , conditional on observing bigram j in quarter t , starting with neutral prior $(1/n_c)$, is:

$$\hat{\rho}_{ijt}^c = \frac{\hat{q}_{ijt}}{\sum_{i \in c} \hat{q}_{ijt}} = \frac{\hat{q}_{ijt}}{n_c \hat{q}_{cjt}}$$

Define partisanship $\hat{\pi}_{it}^c$ as the weighted average of $\hat{\rho}_{ijt}^c$ across bigrams

$$\begin{aligned} \hat{\pi}_{it}^c &= \sum_j \hat{q}_{ijt} \hat{\rho}_{ijt}^c \\ &= \frac{1}{n_c} \sum_j \frac{\hat{q}_{ij}^2}{\hat{q}_{cjt}} \\ &= \frac{1}{n_c} [1 + \hat{\chi}_{ict}^2] \end{aligned}$$

where for the last step we used:

$$\begin{aligned} \hat{\chi}_{ict}^2 &= \sum_j \frac{\hat{q}_{ijt}^2 + \hat{q}_{cjt}^2 - 2\hat{q}_{ijt}\hat{q}_{cjt}}{\hat{q}_{cjt}} \\ &= \sum_j \frac{\hat{q}_{ijt}^2}{\hat{q}_{cjt}} + \sum_j \hat{q}_{cjt} - 2 \sum_j \hat{q}_{ijt} \\ &= \sum_j \frac{\hat{q}_{ijt}^2}{\hat{q}_{cjt}} - 1 \end{aligned}$$

C Choice of fixed covariate penalty ψ

Following GST, the fixed covariate cost should be chosen so that, when the populist penalty is set to an almost negligible level for all bigrams ($\lambda_j = 1e - 10$), the resulting partisanship estimates are similar to those obtained under MLE. As discussed in the main text, the role of λ_j in the main specification

is to reduce the finite-sample bias of the estimator. Therefore, we want the fixed covariate cost ψ calibrated such that, when we do not attempt to reduce finite-sample bias, the estimates resemble the MLE results. At the same time, we prefer to keep the fixed covariate cost relatively high, so that the covariates do not absorb too much of the variation in bigram usage.

To find an optimal value, we explore different fixed costs values – $\psi \in \{1e-03, 1e-04, 1e-05, 1e-06\}$. Both $1e-05$ and $1e-06$ lead to a pattern of partisanship similar to the MLE (results available upon request). We thus choose $1e-05$ as a fixed cost for the rest of the paper.

D Construction of confidence intervals

Given that standard nonparametric bootstrap is known to be invalid for lasso regression (Chatterjee and Lahiri, 2011), confidence intervals are estimated via subsampling, similar to the way it is done in GST. To do so, we randomly draw speakers without replacement to create 500 subsamples each containing (up to integer restrictions) one-fifth of all speakers and, for each subsample k , compute the penalized estimate of partisanship π_t^k . The confidence interval (90%) is then constructed according to the following formula around the estimated π_t^* :

$$CI(\pi_t^*) = \frac{1}{2} + [\exp(\log(\pi_t^* - 1/2) - Q_{t(475)}^{*k}/\sqrt{\tau}); \exp(\log(\pi_t^* - 1/2) - Q_{t(25)}^{*k}/\sqrt{\tau})],$$

where τ is the number of speakers in the full sample, $Q_{t(b)}^{*k}$ is the b th order statistics of:

$$Q_t^{*k} = \sqrt{\tau_k}(\log(\pi_t^k - 1/2) - \log([\frac{1}{100} \sum_{l=1}^{100} \pi_t^l] - 1/2))$$

E International Connections of Populist and Non-populist politicians

As explained in the main text, our emended Coleman index is given by:

$$H_i^{cross} = \frac{f_i - a_i}{1 - a_i}, \quad (11)$$

where f_i is the fraction of outgoing ties to nodes (Twitter accounts of leaders in our sample) who are of the same type as i but in a different country (relative to the total outgoing edges of i), a_i is the fraction of all nodes N that, for i , are of the same type but lie in a different country. When $H_i^{cross} > 0$ the node over-connects with same-type nodes in different countries, while $H_i^{cross} < 0$ means that it under-connects.

Table E1: Cross-Country Coleman Index by Politician Type, 2013-22 average

	P	NP	Difference	P-value
H_i^{cross} (mean)	-0.036	-1.230	1.195***	0.000
H_i^{cross} (weighted mean)	0.017	-0.857	0.874***	0.000

Notes. P-values are from a two-sided permutation test (10,000 replications) of the difference in mean Coleman H index between populist and non-populist politicians. In each permutation, the populist/non-populist labels are randomly reassigned across nodes. The weighted Coleman index is computed by weighting each node's H_i value by its total follows, so actors with more outgoing ties contribute proportionally more to the average.

Table E2: Cross-Country Coleman Index by Politician Type & Political Leaning, 2013-22 average

	Left-Wing			Right-Wing		
	P	NP	Diff.	P	NP	Diff.
H_i^{cross} (mean)	0.009	-0.589	0.598***	0.039	-0.079	0.119***
			$p = 0.000$			$p = 0.003$
H_i^{cross} (weighted mean)	0.052	-0.438	0.490***	0.063	-0.027	0.089**
			$p = 0.000$			$p = 0.021$

Notes. P-values are from a two-sided permutation test (10,000 replications) of the difference in mean Coleman H index between populist and non-populist politicians with the same political leaning. In each permutation, the populist/non-populist labels are randomly reassigned across nodes. The weighted Coleman index is computed by weighting each node's H_i value by its total follows, so actors with more outgoing ties contribute proportionally more to the average.

Table E3: Within-Country Coleman Index by Politician Type, 2013-22 average

	P	NP	Difference	P-value
H_i^{within} (mean)	0.221	0.455	-0.234***	0.000
H_i^{within} (weighted mean)	0.153	0.349	-0.196***	0.000

Notes. P-values are from a two-sided permutation test (10,000 replications) of the difference in mean Coleman H index between populist and non-populist politicians. In each permutation, the populist/non-populist labels are randomly reassigned across nodes. The weighted Coleman index is computed by weighting each node's H_i value by its total follows, so actors with more outgoing ties contribute proportionally more to the average.

Table E4: Within-Country Coleman Index by Politician Type & Political Leaning, 2013-22 average

	Left-Wing			Right-Wing		
	P	NP	Diff.	P	NP	Diff.
H_i^{within} (mean)	0.138	0.369	-0.231*** $p = 0.000$	0.164	0.235	-0.071 $p = 0.224$
H_i^{within} (weighted mean)	0.106	0.295	-0.189*** $p = 0.000$	0.113	0.106	0.007 $p = 0.866$

Notes. P-values are from a two-sided permutation test (10,000 replications) of the difference in mean Coleman H index between populist and non-populist politicians with the same political leaning. In each permutation, the populist/non-populist labels are randomly reassigned across nodes. The weighted Coleman index is computed by weighting each node's H_i value by its total follows, so actors with more outgoing ties contribute proportionally more to the average.

F Classification of Populist Leaders

F.1 Switching political types

We directly asked ChatGPT whether the politician changed between being a populist and a non-populist using the prompt below, and repeated the process 4 times. Over the 367 politicians in our sample, the answer is "Unchanged" for 363 of them with full consistency (4/4). We just have 3 cases where ChatGPT answered 2 times over 4 "Changed from NP to P". These politicians are Cabo Daciolo from Brazil, Mark Latham from Australia, and Juliette Boulet from Belgium. For Paweł Piotr Kukiz 3 out of 4 times ChatGPT gave an answer "Changed from NP to P".

The ChatGPT prompt used to assess whether politicians switched from

one political type to another was as follows:

*I will give you the name of a present or past leader of a political party and the country they are from. I am interested in whether this person ****changed**** between being a populist and a non-populist, or vice-versa, during the period 2010–2022.*

My definition of populism is the following: A leader is defined as populist if he or she divides society into two artificial groups – “the people” vs. “the elites” – and then claims to be the sole representative of the true people. Populists place the alleged struggle of the people (“us”) against the elites (“them”) at the center of their political campaign and governing style. More precisely, populists typically depict “the people” as a suffering, inherently good, virtuous, authentic, ordinary, and common majority, whose collective will is incarnated in the populist leader. By contrast, “the elite” is an inherently corrupt, self-serving, power-hoarding minority, negatively defined as all those who are not “the people”.

*Determine whether the leader ****changed**** between being a populist and a non-populist, or vice-versa, ****between 2010 and 2022**** and give me one of these four possible answers:*

- 1. ****Changed from P to NP****: if the answer is yes and the politician changed from being a Populist to a Non-Populist.*
- 2. ****Changed from NP to P****: if the answer is yes and the politician changed from being a Non-Populist to a Populist.*
- 3. ****Unchanged****: if the answer is no, meaning the politician was either always a Populist or always a Non-Populist.*
- 4. ****Ambiguous****: if the answer is ambiguous because either you have lack of sufficient information or it is truly an ambiguous case.*

Rules: - Output only a single label. - Be concise and accurate.

F.2 Comparison of our classification with fully automated ChatGPT classification

We run the exercise similar to the one above to validate our manual classification of populist politicians. To do so, we asked ChatGPT to classify politicians as either populists, non-populists or ambiguous. The results show a total of 26 cases of disagreement between our classification and the benchmark classification used in the main text (the classification in which the remaining ambiguous cases are considered populists). In all 26 cases, our benchmark classification identifies the politicians as populists, whereas the fully automated ChatGPT classification labels them as non-populists. Of these 26 politicians, 10 are classified as populists in our data because they belong to a populist party, following our approach for ambiguous cases to reduce the overall number of ambiguities.

Among the remaining 16 politicians, 7 are classified as ambiguous and do not belong to a populist party; therefore, following the approach described in the main text, we treat them as populists in the benchmark specification. The other 9 politicians are considered populists in our classification but are labeled as non-populists by the automated ChatGPT classification. These politicians are Beata Maria Kuśińska, Catarina Martins, Elizabeth Warren, Jeremy Corbyn, Judit Varga, Kristian Thulesen Dahl, Pia Kjaersgaard, Roberto Maroni, and Zoltán Kovács.

G Classification of niche parties

To classify niche parties, the following prompt was used:
You are a helpful research assistant. I will give you the name of a past or present political party and what country they are from. I am interested in whether they are a 'niche' political party. In particular, a party is niche if it satisfies all three of the following criteria:

- 1) Niche parties reject the traditional class-based orientation of politics. So instead of prioritizing demands for more or less redistribution, these parties politicize other issues which were previously not represented in party competition. In many cases, green parties or radical right parties would be examples of this.

- 2) The issues raised by niche parties are not only novel, but they often

do not coincide with existing lines of political division. Niche parties appeal to groups of voters that may cross-cut traditional partisan alignments. As a result, cases of voter defection between “unlikely” party pairs have occurred. The defection of former British Conservative voters to the Green Party in 1989 and former French Communist party voters to the radical right Front National in 1986 are typical examples.

3) Niche parties keep a very narrow policy agenda, focusing only on a few key policies. So for example the Green Party in Germany is not a niche party because it is a mainstream party with a very broad policy agenda.

Respond with 'yes' if the party is niche and 'no' if it is not.

H Classification of Leaders along Left-Right

H.1 ChatGPT Classification

Classification of politicians with Chat GPT was performed using the model GPT-4o-mini and with the following prompt:

I will give you the name of a present or past leader of a political party and what country they are from. I am interested in whether they are a left-wing or right-wing politician in the period 2010-2021. In particular, give me one of these three possible answers:

1. left: commonly considered a left-wing leader 2. right: commonly considered a right-wing leader 3. ambiguous: ambiguous whether they are left- or right-wing because either: (i) they are truly an ambiguous case (ii) they change from being left- to right-wing (or vice-versa) during the period (iii) you do not have enough information about the individual to give an answer

Rules: - Output only a single label. - If the label is ambiguous, specify whether it is because of (i), (ii) or (iii). - Be concise and accurate.

ChatGPT classification exhibits some randomness, therefore results slightly change when the same procedure is repeated. Positively, in the 4 iterations made there was no case in which Chat GPT switched from classifying a politician from right to left or vice-versa, but just observations where the classifica-

tion switched between right and ambiguous or between left and ambiguous.

To account for this, the classification has been repeated 4 times, and politicians have been classified as "Right" if ChatGPT classified the politician in this category all 4 times, the same with "Left". Otherwise, politicians have been assigned to the "ambiguous" category. This process led to 82 ambiguous cases (later 64 of these politicians were classified as left by the RILE classification, 18 as right). The number of politicians which were positively classified is 285: the confusion matrices on this sample with the RILE classification are reported below – Table H1.

H.2 RILE Classification

We also classify all the politicians using the party-level classification. The Manifesto Project Dataset measures left-right positions of political parties through the ‘rile’ index, which measures how much a party talks about left-wing or right-wing issues. We consider, for each party, the most recent codification of this index and classify politicians as right if they belong to a party whose variable is greater than zero, as left otherwise. Politicians that have no match with the dataset have been manually classified (with the use of ChatGPT).

We prefer the ChatGPT classification to the one based on the Manifesto Project Dataset, as we believe that the former could better capture within-party factions which are relevant in classifying individual leaders, and it is more comparable internationally. In the Appendix Figure L22 we present our main results for the L-R classification based on RILE index instead of ChatGPT.

Table H1: Left-Right Classifications Confusion Matrix - Absolute Values

	GPT Left	GPT Right	
RILE Left	147	46	193
RILE Right	12	80	92
	159	126	285

I Summary of leaders classifications

Australia

Andrew Bartlett (NP, L), Bill Shorten (NP, L), Bob Katter (P, R), Campbell Newman (NP, R), Christine Milne (NP, L), Clive Palmer (P, R), Deb Frecklington (NP, R), John Anderson (NP, L), John-Paul Langbroek (NP, R), Julia Gillard (NP, L), Kevin Rudd (NP, L), Kim Beazley (NP, L), Malcolm Turnbull (NP, R), Mark Latham (P, R), Natasha Stott Despoja (NP, L), Nick Xenophon (P, L), Pauline Hanson (P, R), Richard Di Natale (NP, L), Scott Morrison (NP, R), Tim Nicholls (NP, R), Tony Abbott (P, R)

Austria

Alexander Van der Bellen (NP, L), Beate Meinl-Reisinger (NP, L), Christian Kern (NP, L), Heinz-Christian Strache (P, R), Josef Bucher (P, R), Matthias Strolz (NP, L), Norbert Hofer (P, R), Pamela Rendi-Wagner (NP, L), Peter Pilz (NP, L), Sebastian Kurz (P, R), Ulrike Lunacek (NP, L), Werner Kogler (NP, L)

Belgium

Alexander De Croo (NP, R), Bart De Wever (P, R), Benoît Lutgen (NP, R), Caroline Gennez (NP, L), Charles Michel (NP, R), Didier Reynders (NP, L), Elio Di Rupo (NP, L), Guy Verhofstadt (NP, L), Gwendolyn Rutten (NP, R), Herman Van Rompuy (NP, L), Jean-Marc Nollet (NP, L), John Crombez (NP, L), Joëlle Milquet (NP, R), Juliette Boulet (NP, L), Marianne Thyssen (NP, R), Meyrem Almaci (NP, L), Peter Mertens (P, L), Philippe Henry (NP, L), Philippe Lamberts (NP, L), Stefaan De Clerck (NP, R), Tom Van Grieken (P, R), Wouter Beke (NP, R), Wouter Van Besien (NP, L), Yves Leterme (NP, L), Zoé Genot (NP, L)

Brazil

Anthony Garotinho (P, L), Aécio Neves (NP, R), Cabo Daciolo (P, R), Ciro Gomes (P, L), Cristovam Buarque (NP, L), Dilma Rousseff (P, L), Fernando Haddad (NP, L), Geraldo Alckmin (NP, R), Heloísa Helena (P, L), Henrique Meirelles (NP, R), Jair Bolsonaro (P, R), José Serra (NP, R), João Amoêdo (NP, R), Luciana Genro (NP, L), Luiz Inácio Lula da Silva (P, L), Marina Silva (NP, L), Rui Costa Pimenta (NP, L)

Czech Republic

Andrej Babiš (P, L), Ivan Bartoš (NP, L), Jan Farský (NP, L), Jiří Paroubek (NP, L), Jiří Rusnok (NP, R), Karel Schwarzenberg (NP, L), Lubomír Zaorálek (NP, L), Mirek Topolánek (NP, R), Miroslav Kalousek (NP, R), Miroslava Němcová (NP, R), Pavel Bělobrádek (NP, R), Petr Fiala (NP, R), Tomio Okamura (P, R), Vladimír Špidla (NP, L), Vojtěch Filip (NP, L)

Denmark

Anders Fogh Rasmussen (NP, R), Anders Samuelsen (NP, R), Bendt Bendtsen (NP, R), Helle Thorning-Schmidt (NP, L), Holger K. Nielsen (NP, L), Kristian Thulesen Dahl (P, R), Lars Barfoed (NP, R), Lars Løkke Rasmussen (NP, L), Margrethe Vestager (NP, L), Marianne Jelved (NP, L), Mogens Lykketoft (NP, L), Morten Østergaard (NP, L), Pernille Skipper (P, L), Pernille Vermund (P, R), Pia Kjaersgaard (P, R), Pia Olsen Dyhr (NP, L), Søren Pape Poulsen (NP, R), Uffe Elbæk (NP, L)

Finland

Alexander Stubb (NP, R), Anna-Maja Henriksson (NP, L), Anneli Jäätteenmäki (NP, L), Anni Sinnemäki (NP, L), Antti Rinne (NP, L), Carl Haglund (NP, L), Eero Heinäluoma (NP, L), Harry Harkimo (NP, L), Jussi Halla-aho (P, R), Jutta Urpilainen (NP, L), Jyrki Katainen (NP, R), Li Andersson (NP, L), Mari Kiviniemi (NP, R), Matti Vanhanen (NP, R), Osmo Soininvaara (NP, L), Paavo Arhinmäki (NP, L), Pekka Haavisto (NP, L), Petteri Orpo (NP, R), Päivi Räsänen (NP, R), Sari Essayah (NP, R), Suvi-Anne Siimes (NP, L), Tarja Cronberg (NP, L), Ville Niinistö (NP, L)

France

Benoît Hamon (NP, L), Bernard Cazeneuve (NP, L), Dominique Voynet (NP, L), Dominique de Villepin (NP, R), Edouard Philippe (NP, R), Emmanuel Macron (NP, L), Eva Joly (NP, L), François Bayrou (NP, L), François Fillon (NP, R), François Hollande (NP, L), Jean-Luc Mélenchon (P, L), Jean-Marc Ayrault (NP, L), Jean-Marie Le Pen (P, R), Jean-Pierre Raffarin (NP, R), Manuel Valls (NP, L), Marine Le Pen (P, R), Nicolas Dupont-Aignan (P, R), Nicolas Sarkozy (NP, R), Noël Mamère (NP, L), Ségolène Royal (NP, L)

Germany

Alice Weidel (P, R), Bernd Lucke (P, R), Cem Özdemir (NP, L), Christian Lindner (NP, R), Dietmar Bartsch (NP, L), Gabi Zimmer (NP, L), Gregor Gysi (NP, L), Jürgen Trittin (NP, L), Katrin Göring-Eckardt (NP, L), Martin Schulz (NP, L), Peer Steinbrück (NP, L), Renate Künast (NP, L), Sahra Wagenknecht (P, L)

Hungary

Bocskor Andrea (NP, R), Deli Andor (NP, L), Deutsch Tamás (P, R), Edina Toth (P, R), Gergely Karácsony (NP, L), Gordon Bajnai (NP, L), István Ujhelyi (NP, L), Judit Varga (P, R), Katalin Novák (P, R), László Andor (NP, L), Márton Gyöngyösi (P, R), Péter Niedermüller (NP, L), Tibor Navracsic (NP, R), Zoltan Kovacs (P, R)

Italy

Alessandro Di Battista (P, L), Beppe Grillo (P, L), Carlo Calenda (NP, L), Daniela Santanché (P, R), Emma Bonino (NP, L), Enrico Letta (NP, L), Francesco Rutelli (NP, L), Giorgia Meloni (P, R), Giuseppe Conte (P, L), Luigi Di Maio (P, L), Mario Monti (NP, R), Matteo Renzi (NP, L), Matteo Salvini (P, R), Paolo Gentiloni (NP, L), Pier Ferdinando Casini (NP, R), Pier Luigi Bersani (NP, L), Pietro Grasso (NP, L), Roberto Maroni (P, R), Silvio Berlusconi (P, R), Walter Veltroni (NP, L)

Mexico

Alberto Anaya Gutiérrez (NP, L), Alejandro Moreno (NP, R), Alfonso Ramírez Cuéllar (P, L), Andrés Manuel López Obrador (P, L), Beatriz Mojica Morga (NP, L), Carlos Puente Salas (NP, L), Carolina Viggiano (NP, R), Clemente Castañeda Hoefflich (NP, L), Dante Delgado (NP, R), Enrique Peña Nieto (NP, R), Felipe Calderón (NP, R), Gabriel Quadri de la Torre (NP, R), Hugo Eric Flores Cervantes (P, L), Héctor Larios (NP, R), Jaime Rodríguez Calderón (P, L), Josefina Vázquez Mota (NP, R), José Antonio Meade (NP, L), Luis Castro Obregón (NP, L), Manuel Granados Covarrubias (NP, L), Marko Cortés (NP, R), Martí Batres (P, L), Patricia Mercado Castro (NP, L), Reginaldo Sandoval Flores (P, L), Ricardo Anaya (NP, R), Roberto Campa Cifrián (NP, L), Yeidckol Polevnsky (P, L)

Netherlands

Alexander Pechtold (NP, L), André Rouvoet (NP, L), Arie Slob (NP, R), Diederik Samsom (NP, L), Emile Roemer (P, L), Geert Wilders (P, R), Gert-Jan Segers (NP, R), Henk Krol (NP, L), Jan Peter Balkenende (NP, R), Jesse Klaver (NP, L), Job Cohen (NP, L), Kees van der Staaij (NP, R), Lodewijk Asscher (NP, L), Marianne Thieme (NP, L), Mark Rutte (NP, R), Paul Rosenmöller (NP, L), Sybrand van Haersma Buma (NP, R), Thierry Baudet (P, R), Thom de Graaf (NP, L), Tunahan Kuzu (P, L)

Norway

Audun Lysbakken (NP, L), Bjørnar Moxnes (NP, L), Dagfinn Høybråten (NP, R), Erna Solberg (NP, R), Hanna E. Marcussen (NP, L), Jens Stoltenberg (NP, L), Jonas Gahr Støre (NP, L), Knut Arild Hareide (NP, L), Kristin Halvorsen (NP, L), Liv Signe Navarsete (P, L), Rasmus Hansson (NP, L), Thorbjørn Jagland (NP, L), Trine Skei Grande (NP, R), Une Aina Bastholm (NP, L), Åslaug Haga (P, L)

Poland

Barbara Nowacka (NP, L), Beata Maria Kusińska (P, R), Donald Tusk (NP, L), Ewa Kopacz (NP, L),

Grzegorz Napieralski (NP, L), Janusz Korwin-Mikke (P, R), Janusz Palikot (P, R), Janusz Piechociński (NP, L), Leszek Miller (NP, L), Mateusz Morawiecki (P, R), Małgorzata Maria Kidawa-Błońska (NP, L), Paweł Piotr Kukiz (P, L), Roman Giertych (NP, R), Ryszard Petru (NP, R), Wojciech Olejniczak (NP, L), Władysław Marcin Kosiniak-Kamysz (NP, L), Włodzimierz Czarzasty (NP, L)

Portugal

André Ventura (P, R), António Costa (NP, L), Assunção Cristas (NP, R), Carlos Guimarães Pinto (NP, R), Catarina Martins (P, L), José Manuel Barroso (NP, L), José Sócrates (NP, L), Mário Centeno (NP, L), Pedro Passos Coelho (NP, R), Rui Rio (NP, R)

Slovenia

Alenka Bratušek (NP, L), Borut Pahor (NP, L), Dejan Židan (NP, L), Franc Bogovič (NP, R), Gregor Golobič (NP, L), Gregor Virant (NP, L), Janez Janša (P, R), Janez Podobnik (NP, R), Karl Erjavec (NP, L), Katarina Kresal (NP, L), Ljudmila Novak (NP, R), Luka Mesec (P, L), Marjan Šarec (P, L), Matej Tonin (NP, R), Miro Cerar (NP, L), Zmago Jelinčič Plemeniti (P, R), Zoran Janković (NP, R)

Spain

Albert Rivera (NP, R), Cayo Lara (NP, L), Francesc Homs (NP, L), Gabriel Rufián (P, L), Gaspar Llamazares (NP, L), Joan Ridao (P, L), Josu Erkoreka (NP, L), Mariano Rajoy Brey (NP, R), Oriol Junqueras (P, L), Pablo Casado (NP, R), Pablo Iglesias (P, L), Pedro Sánchez (NP, L), Rosa Díez (NP, L), Santiago Abascal (P, R)

Sweden

Alf Svensson (NP, L), Annie Lööf (NP, R), Ebba Busch Thor (P, R), Gudrun Schyman (NP, L), Göran Hägglund (NP, R), Göran Persson (NP, L), Isabella Lövin (NP, L), Jan Björklund (NP, R), Jimmie Åkesson (P, R), Jonas Sjöstedt (NP, L), Lars Ohly (NP, L), Maria Wetterstrand (NP, L), Stefan Löfven (NP, L), Åsa Romson (NP, L)

United Kingdom

Arlene Foster (NP, R), Boris Johnson (P, R), Charles Kennedy (NP, L), David Cameron (NP, R), Ed Miliband (NP, L), Gordon Brown (NP, L), Jeremy Corbyn (P, L), Jo Swinson (NP, L), John Swinney (NP, L), Nicola Sturgeon (NP, L), Nigel Farage (P, R), Paddy Ashdown (NP, L), Theresa May (NP, R), Tim Farron (NP, L), William Hague (NP, R)

United States

Alexandria Ocasio-Cortez (P, L), Barack Obama (NP, L), Bernie Sanders (P, L), Bill Frist (NP, R), Chuck Schumer (NP, L), Elizabeth Warren (P, L), George Bush (NP, R), Harry Reid (NP, L), Hillary Clinton (NP, L), Joe Biden (NP, L), John Boehner (NP, R), John Kerry (NP, L), John McCain (NP, R), José E. Serrano (NP, L), Kevin McCarthy (NP, R), Mike Pence (NP, R), Mitch McConnell (NP, R), Mitt Romney (NP, R), Nancy Pelosi (NP, L), Paul Ryan (NP, R), Steny Hoyer (NP, L)

J Bigrams selection

Figure J1 restricts the visualization to bigrams that appear in at least two quarters (left panel) and that occur at least ten times in at least one quarter (right panel). The Figure guides us in our choice of bigram-selection thresholds. The benchmark specification uses bigrams used in at least 10 quarters with a minimum of 25 occurrences in at least one quarter. For robustness

checks we also focus on bigrams that were used in at least 5/10 quarters and at least 50 times in at least one of the quarters.

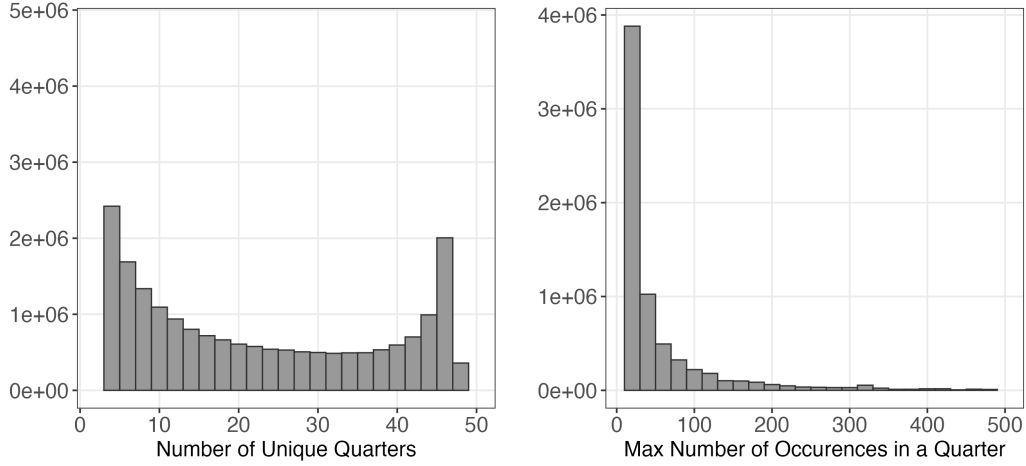


Figure J1: Histograms for the two measures of bigrams frequency: No. of unique quarters a bigram was used (left panel), Max no. of occurrences of a bigram in a quarter (right panel). Low values of both measures are omitted.

K Monte Carlo simulation for number of speakers and unobserved groups

Figure 2 raises two concerns: (1) partisanship for the random group exceeds 0.5, (2) partisanship for the random group of politicians exhibit pre-electoral divergence (qualitatively similar to the dynamics for the populist divide). Here we conduct a Monte Carlo simulation to examine whether these patterns could be due to the small number of politicians in the sample (relative to the number of possibly relevant political types).

We find that the observed patterns for the random groups are likely to be driven by both (1) and (2). However, it is important to note that the results obtained in this section reflects the properties of the penalized estimators when being used on different samples, not the properties of the samples. We verify that this is indeed the case by estimating random groups partisanship with the leave-out estimator, which consistently yields partisanship of 0.5 even in the presence of unobserved groups.

We keep the same set of bigrams used in the benchmark specification ($N = 6144$), but generate random bigram usage along with random samples of politicians for 20 sessions. For the politician samples, we construct a

populist dummy, preserving the observed share of populist politicians (approximately 0.225), and vary the number of politicians per sample across $N \in \{50, 250, 500\}$. The median number of active politicians in our sample is 287, so the sample of 250 politicians is the closest to our data configuration. For the bigrams sample, we randomly draw bigram usage using Dirichlet sampling, keeping verbosity close to its observed value. We try to keep the distribution of bigrams such that partisanship matches the one estimated using the real data. Having obtained the simulated sample of bigrams used by each politician: (i) we estimate the parameters of the utility of speech; (ii) compute $\hat{q}_{ijt}^{P_i}$, for each politician i bigram j and session t ; (iii) obtain individual - session specific partisanship measure, π_{it} ; (iv) compute average partisanship, π_t , over the entire sample of politicians.

The bigrams sample is simulated under two scenarios, which we now describe, to explore the importance of small number of politicians and of unobserved political types.

Scenario I. Here we explore the consequences of changing the number of politicians in the sample. First, we randomly select 1,000 bigrams and classify them as populist. When generating bigram usage, we boost the frequency of these bigrams for *populist* politicians by a factor of two. Second, among the 20 sessions we simulate, we treat the first five as pre-electoral, i.e., sessions in which the effect of political types on bigram choice (parameter φ) is higher. To capture this, we apply an additional boost to the same set of 1,000 bigrams in the first five sessions, of about 20% on average. In other words, populist politicians consistently use this set of 1,000 bigrams more frequently across all 20 sessions, but especially so during the first five.

The results of this simulation are presented in Figure K1. The solid lines show the average partisanship outcome, π_t , for the populist versus non-populist division, while the dotted lines depict the average π_t for ten random binary partitions that do not correspond to the populist–non-populist split. Vertical segments are 95% confidence intervals.

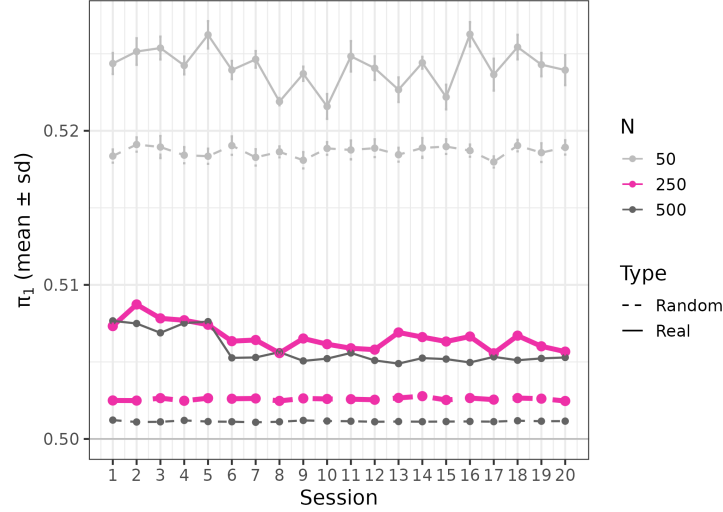


Figure K1: Average partisanship estimates over politicians and sessions for real and random populist groups. No unobserved electorally relevant groups.

As the number of politicians increases, populist partisanship is not affected once it reaches 250 politicians, but partisanship of the random groups (the dotted lines) converges toward 0.5. This addresses concern (1), suggesting that the relatively large observed values of partisanship for the random groups are partly driven by the small number of politicians in our sample.

Nevertheless, in these simulations partisanship of the random groups is not sensitive to varying parameter φ , since it has the same pattern in the first five sessions as in the remaining ones (the difference between the first sessions and the rest is statistically significant for populist partisanship, but not for partisanship of the random groups). This suggests that the number of politicians, on its own, cannot explain the observed pre-electoral pattern of the random group partisanship in Figure 5.

Scenario II. To explain this pre-electoral pattern, in Figure K2 we explore the consequences of also having unobserved political types. Here we keep the same data generating process, except that in addition to populist and non-populist, we generate two additional political types within the populist and non-populist subsets, boosting for each type the usage of a specific set of bigrams. Moreover, similar to the additional boost in session 1-5 for the populist bigrams, we boost the usage of these specific sets of bigrams in sessions 1-5 too, thus making these four political groups electorally relevant. We then compute average partisanship between two groups of politicians: populist vs

non-populist (the solid lines) and an entirely random partition (the dotted lines). Again we repeat these simulations for three sample sizes of politicians, 50, 250 and 500. The results of these simulations are depicted in Figure K2

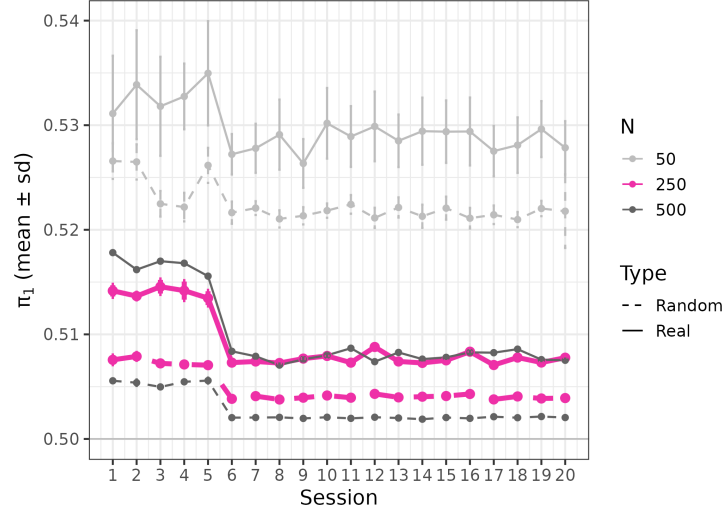


Figure K2: Average partisanship estimates over politicians and sessions for real and random populist groups. Four unobserved electorally relevant groups.

Partisanship in the random groups now is indeed much larger in Sessions 1-5, where unobserved political types have a larger effect on bigram usage, than in the remaining sessions, and this difference is statistically significant. This pattern suggests that some political types are not fully accounted for in our framework. Note that the sample size of politicians continues to be relevant also in these simulations.

This suggests that what matters is not merely the number of politicians, but the number of politicians relative to the number of relevant political types. If the number of politicians is not sufficiently large relative to the number of relevant political types, the two random groups differ systematically in their bigram usage and this is captured by the partisanship statistics.

L Tables and Figures

L.1 Tables

Table L1: Politicians by ideology and populism (counts with row percentages)

	Non-populist	Populist	Total
Left	185 (83.0%)	38 (17.0%)	223 (60.8%)
Right	96 (66.7%)	48 (33.3%)	144 (39.2%)
Total	281 (76.6%)	86 (23.4%)	367 (100%)

Table L2: Wald Tests: Coefficient differences from $\beta_{\text{quarter diff}=-1}$ and $\beta_{\text{quarter diff}=0}$

Hypothesis	F-statistic	p-value	Dep. Variable
$\beta_{\text{quarter diff}=1} = \beta_{\text{quarter diff}=-1}$	3.278	0.07390	χ_{it}
$\beta_{\text{quarter diff}=2} = \beta_{\text{quarter diff}=-1}$	0.201	0.65500	χ_{it}
$\beta_{\text{quarter diff}=3} = \beta_{\text{quarter diff}=-1}$	0.398	0.53000	χ_{it}
$\beta_{\text{quarter diff}=1} = \beta_{\text{quarter diff}=0}$	7.809	0.00649	χ_{it}
$\beta_{\text{quarter diff}=2} = \beta_{\text{quarter diff}=0}$	1.295	0.25900	χ_{it}
$\beta_{\text{quarter diff}=3} = \beta_{\text{quarter diff}=0}$	1.355	0.24800	χ_{it}
$\beta_{\text{quarter diff}=1} = \beta_{\text{quarter diff}=-1}$	5.606	0.02810	π_{it}
$\beta_{\text{quarter diff}=2} = \beta_{\text{quarter diff}=-1}$	7.297	0.01370	π_{it}
$\beta_{\text{quarter diff}=3} = \beta_{\text{quarter diff}=-1}$	3.396	0.08020	π_{it}
$\beta_{\text{quarter diff}=1} = \beta_{\text{quarter diff}=0}$	5.987	0.02380	π_{it}
$\beta_{\text{quarter diff}=2} = \beta_{\text{quarter diff}=0}$	6.886	0.01630	π_{it}
$\beta_{\text{quarter diff}=3} = \beta_{\text{quarter diff}=0}$	1.951	0.17800	π_{it}

Table L3: Unique Elections by Quarter

No. of unique countries per election quarter	1	2	3	4	5
No. of occurrences	12	10	3	2	2

Notes. Only calendar quarter used in the main analysis are considered. Elections that fall outside of the study window are not considered.

Table L4: BERT vs Manual Labeling Metrics

	Accuracy	F1 Score
<i>Policy Topics</i>		
Business	0.91	0.64
Climate	0.96	0.83
Energy	0.97	0.87
European Union	0.98	0.89
Foreign Policy	0.92	0.61
Health	0.97	0.86
Immigration	0.98	0.88
Inequality	0.92	0.55
Inflation	0.96	0.55
Labor Market	0.96	0.83
Rights	0.93	0.68
Security	0.94	0.69
Taxes	0.97	0.86
Trade	0.96	0.42
<i>Non-Policy Topics</i>		
Anti Establishment	0.94	0.49
Personal Communication	0.96	0.53
Elections	0.96	0.67
Self Promotion	0.81	0.33

Notes: Accuracy is defined as the number of correct predictions divided by the total number of predictions. The F1 Score is the harmonic mean of precision and recall: it is twice their product divided by their sum. Precision is the number of true positives divided by the number of all predicted positives, while recall is the number of true positives divided by the number of actual positives.

L.2 Figures

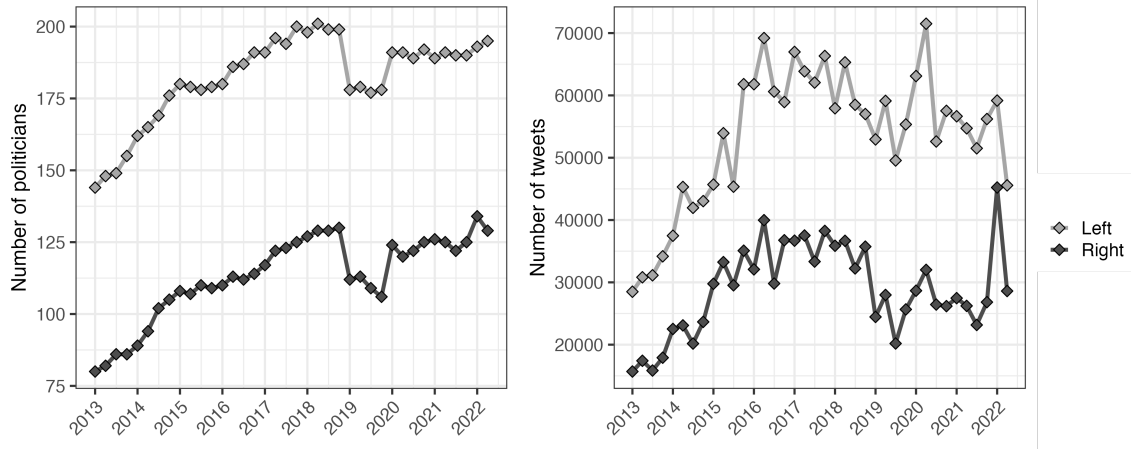


Figure L1: Number of politicians active on Twitter per calendar quarter, and number of tweets, per Left-Right divide

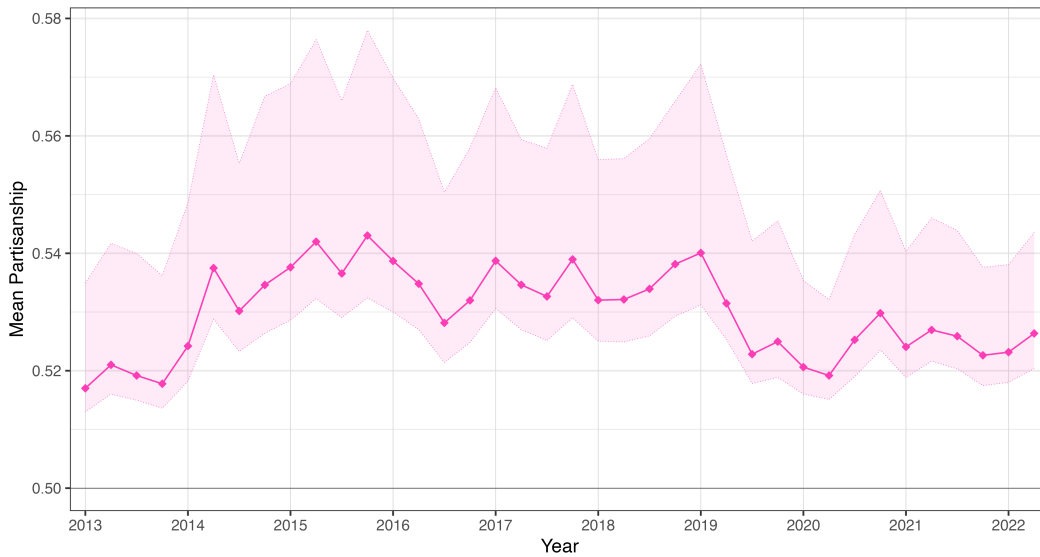


Figure L2: Average partisanship with 90% CI

Notes. Figure plots average values and confidence intervals of populist partisanship over the study period. Partisanship is defined as in Equation 5. Average values are defined as in Equation 6. Confidence Intervals are constructed via subsampling, as described in Section D.

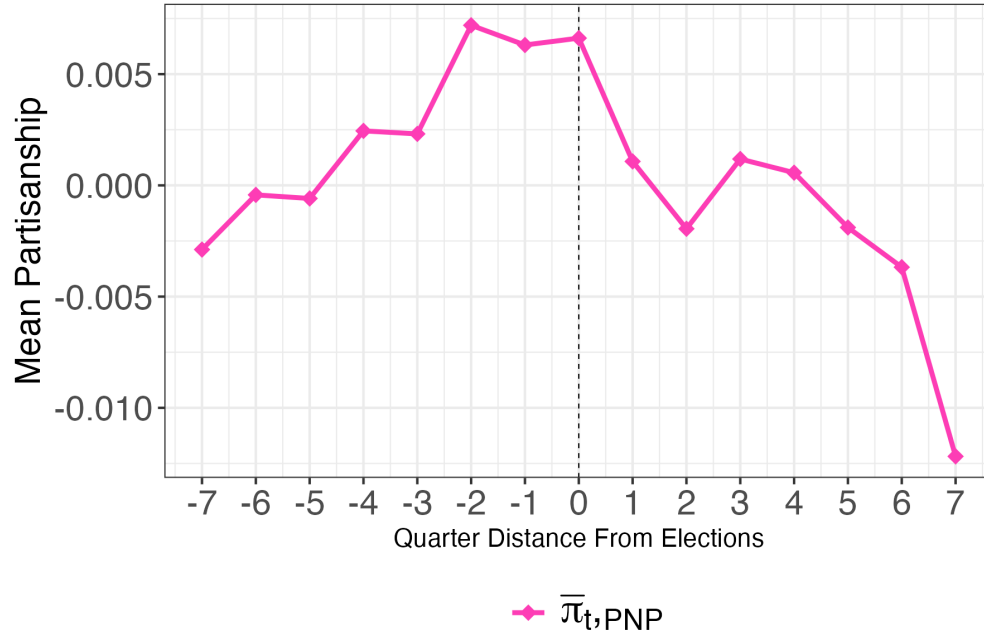


Figure L3: Average estimated π_{it} per quarter difference from elections (net of calendar-quarter fixed effects).

Notes. $\bar{\pi}_{t,PNP}$ refers to average populist partisanship net of calendar-quarter fixed effects. Partisanship is defined as in Equation 5. Average values are defined as in Equation 6.

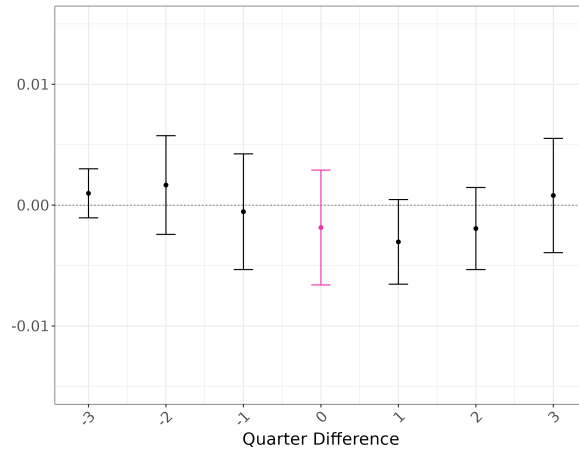


Figure L4: Event studies for π_{it} , Left-Right partisanship measure. Errors are clustered at the elections level, 95% CI are reported.

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from Specification 9. Outcome variable is left-right partisanship. Errors are clustered at the elections level, 95% CI are reported. Partisanship is defined as in Equation 5.

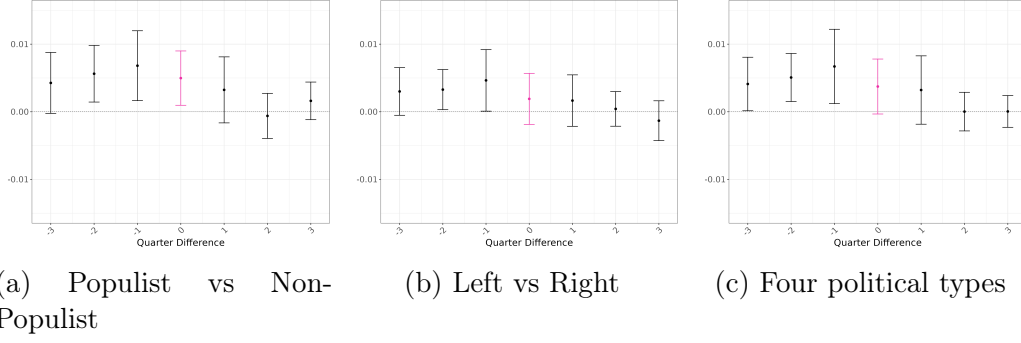


Figure L5: Event study estimates from the model with four political types (include both populist-non-populist and left-right dimensions)

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from Specification 9. Outcome variable is populist partisanship (left panel), left-right partisanship (middle panel) and four types partisanship (right panel). Errors are clustered at the elections level, 95% confidence intervals are reported. Partisanship with several political types is defined as in Subsection A.

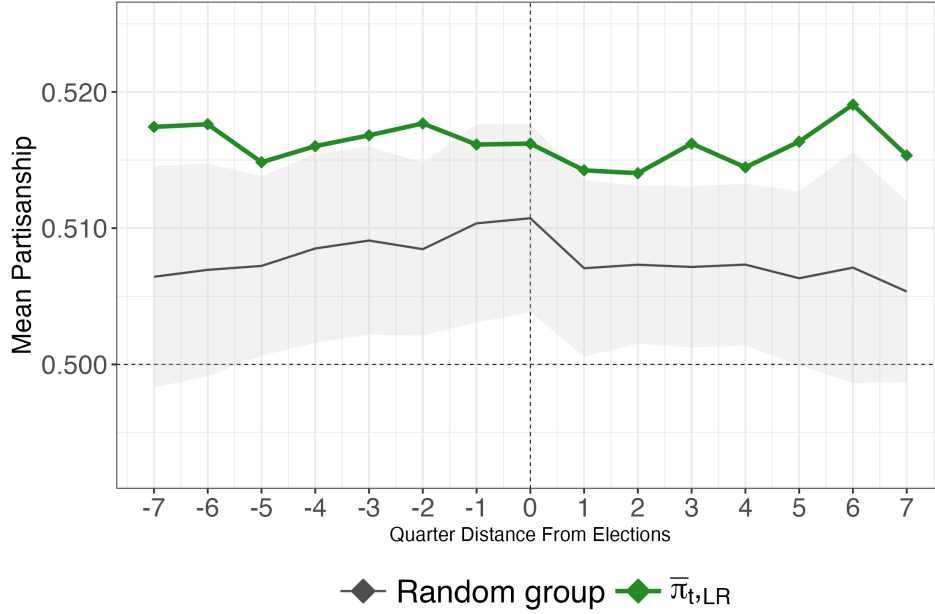


Figure L6: Average estimated χ_{it}^2 and π_{it} per quarter difference from the elections.

Notes. $\bar{\pi}_{t,LR}$ refers to average partisanship defined in the model with two political types: left and right; Random group refers to average partisanship with two randomly defined political types; the shaded area behind the random group is the 95% confidence interval. Average values are defined as in Equation 6.

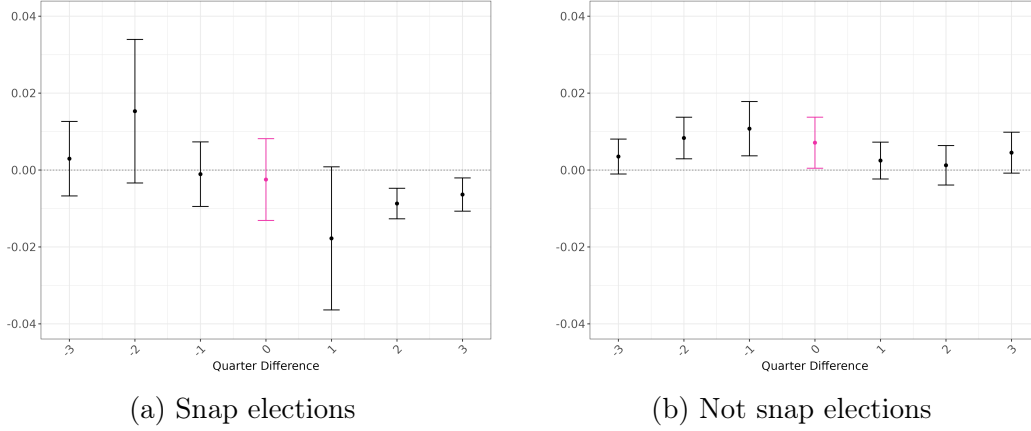


Figure L7: Event study estimates for π_{it} for snap elections (left panel) and elections held in standard time (right panel)

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from Specification 9 on the subsample of snap elections (left panel) and elections held at a regular time (right panel). Errors are clustered at the elections level, 95% confidence intervals are reported. Partisanship is defined as in Equation 5. Snap elections have been manually classified, with the use of both ChatGPT and Wikipedia. We consider elections to be “snap” if they are held at least 6 months before the natural end of the legislation. Out of 97 unique elections in our sample, 20 elections were classified as “snap”.

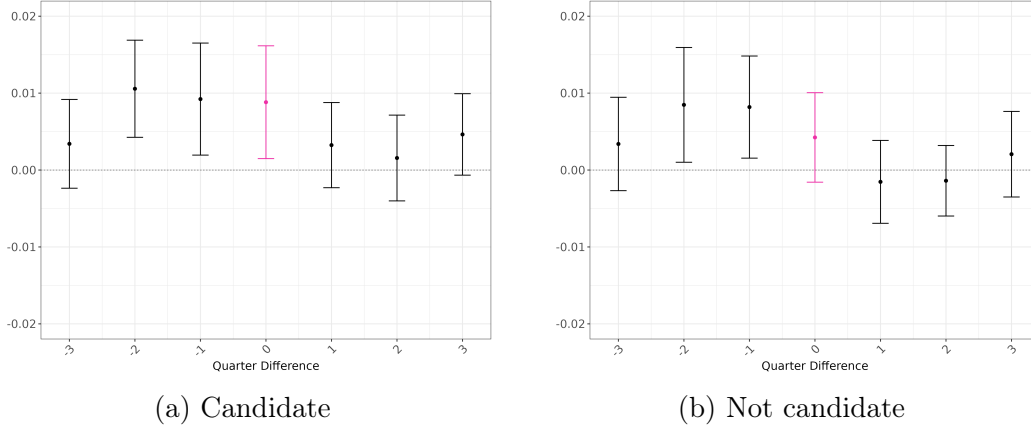


Figure L8: Event study estimates for π_{it} separately for politicians running for reelection in given elections (candidate, left panel), and those who are not running (not candidate, right panel)

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from Specification 9 on the subsample of politicians running for office in a given electoral cycle (left) and those not running (right). Errors are clustered at the elections level, 95% confidence intervals are reported. Partisanship is defined as in Equation 5.

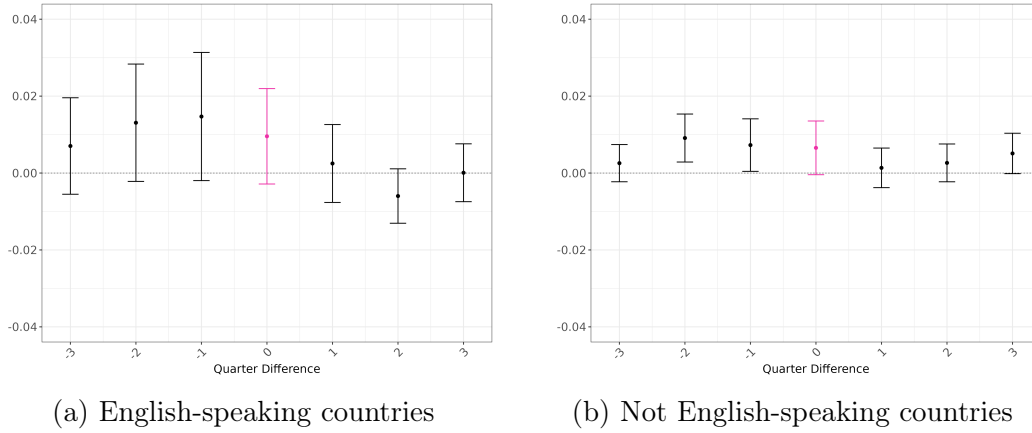


Figure L9: Event study estimates for π_{it} for countries with English and other languages as main official language

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from Specification 9 on the subsamples of English-speaking countries (left) and not English-speaking countries (right). Errors are clustered at the elections level, 95% confidence intervals are reported. Partisanship is defined as in Equation 5.

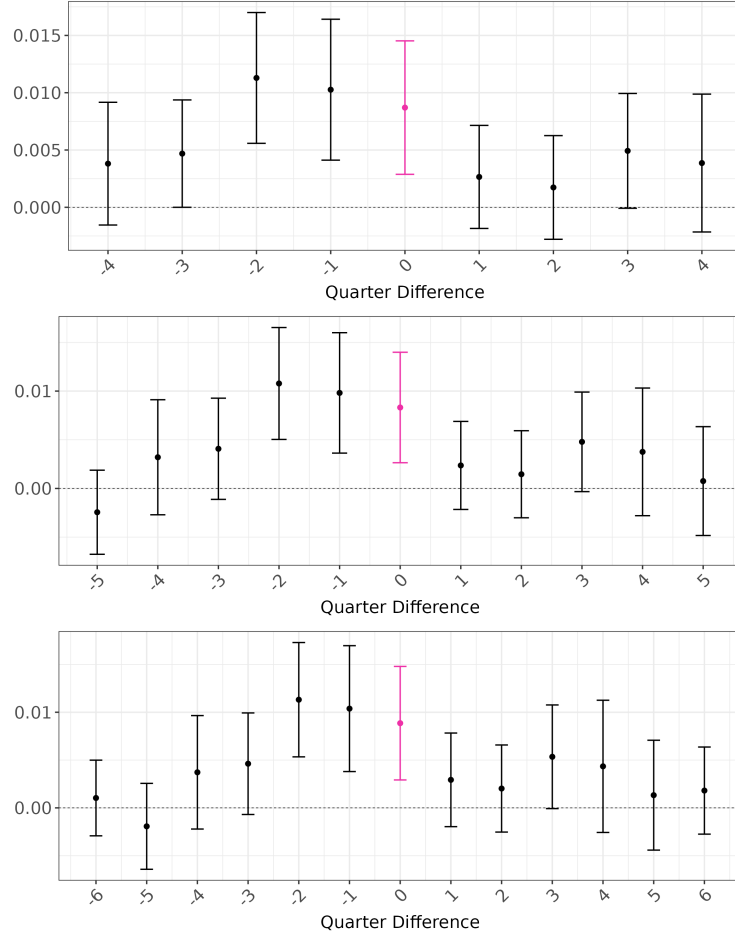


Figure L10: Event study results. Data chunk used is 10-25. Event window is 4 quarters, 5 quarters and 6 quarters.

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from Specification 9 using different definitions of electoral window – 4 quarters (upper panel), 5 quarters (middle panel), 6 quarters (bottom panel). Errors are clustered at the elections level, 95% Confidence intervals are reported. Outcome variable is a populist partisanship, defined as in Equation 5.

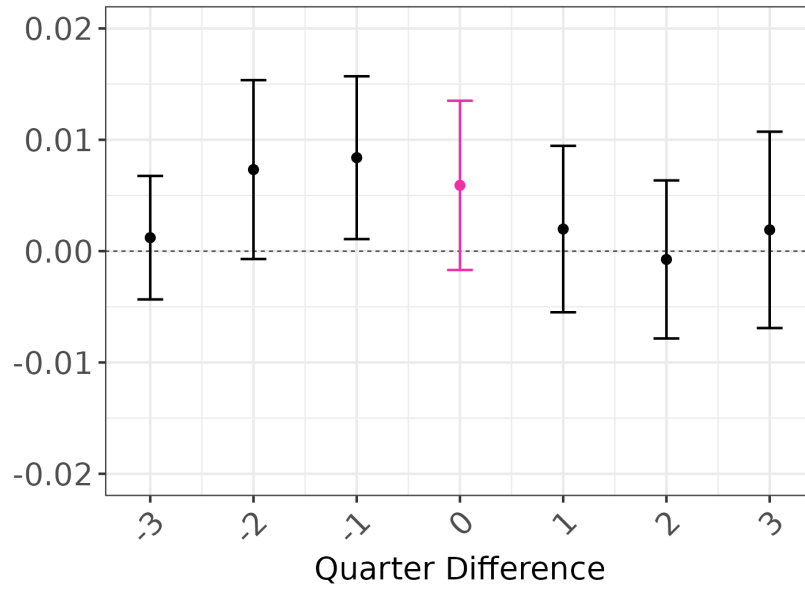


Figure L11: Event study results, staggered difference-in-difference estimations a la Sun and Abraham

Notes. Electoral cycle constitutes a unit of observation. Outcome variable is populist partisanship. Errors are clustered at the elections level, 95% confidence intervals are reported. Partisanship is defined as in Equation 5.

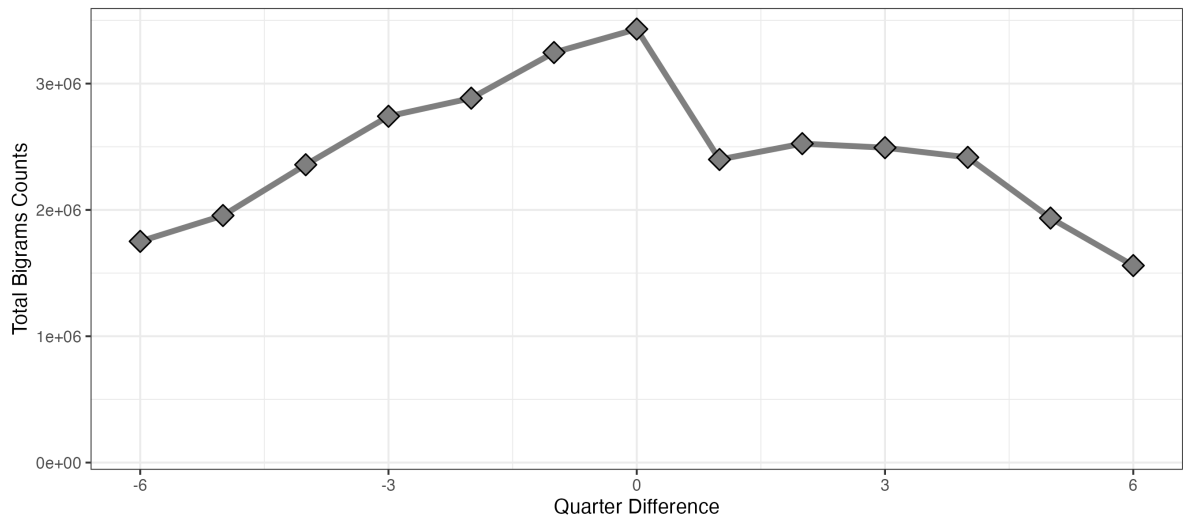


Figure L12: Number of bigrams used per quarter, difference from the elections

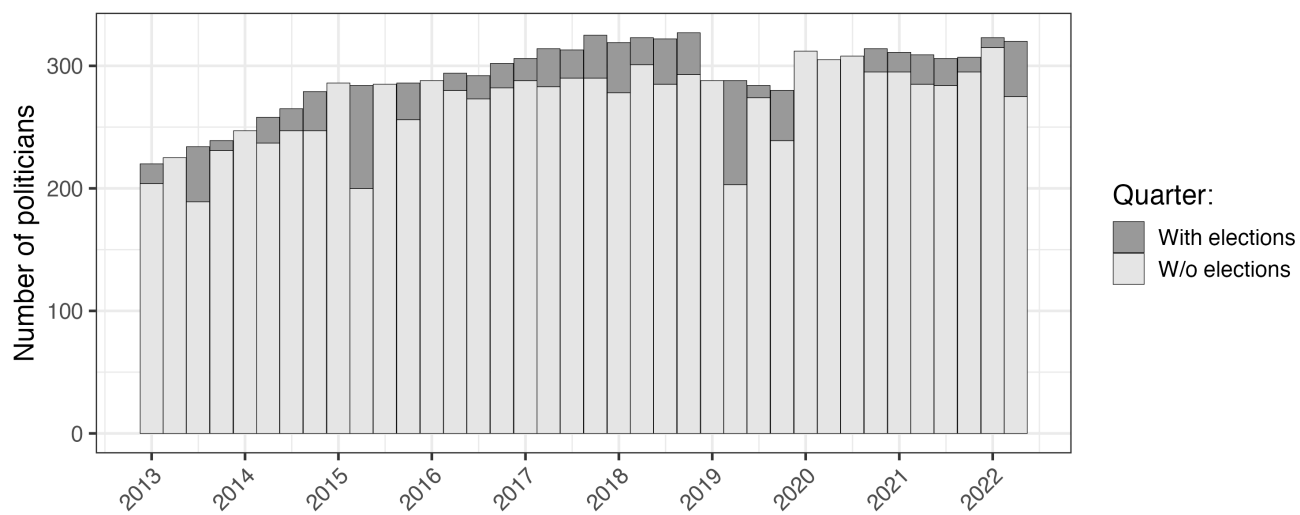


Figure L13: Number of speakers per calendar quarter

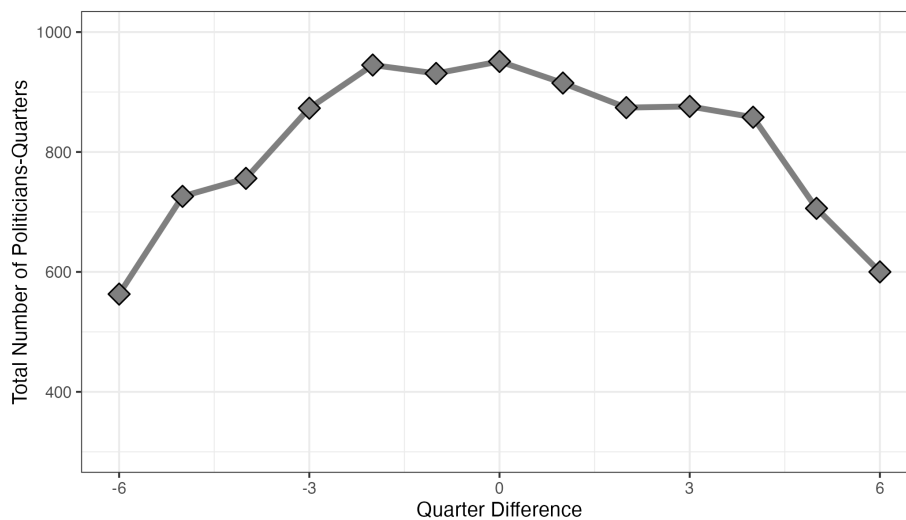


Figure L14: Number of unique quarter-politicians observations across the electoral cycle

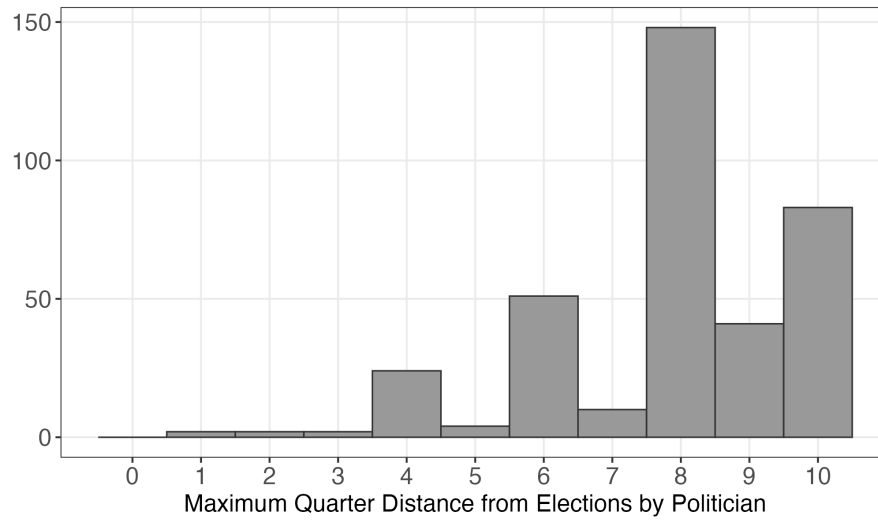


Figure L15: Maximum absolute quarter distance from the elections for politicians

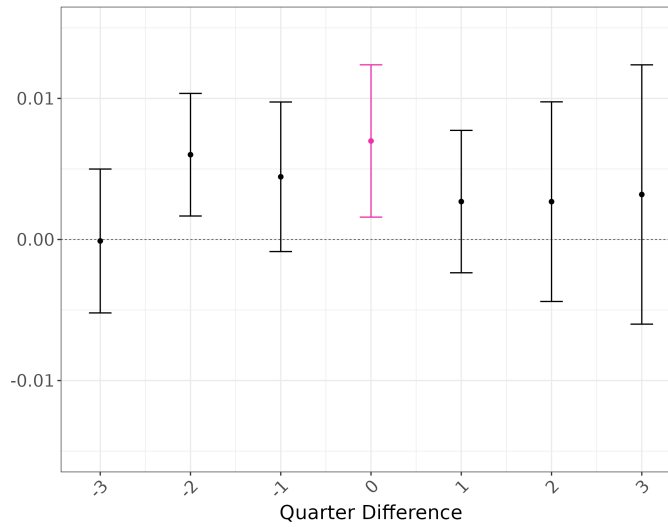


Figure L16: Event study estimates for π_{it} based on a restricted random sample of 220 politicians per quarter.

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from Specification 9 on the restricted sample of 220 politicians randomly selected in each quarter. Errors are clustered at the elections level, 95% confidence intervals are reported. Outcome variable is populist partisanship, defined as in Equation 5.

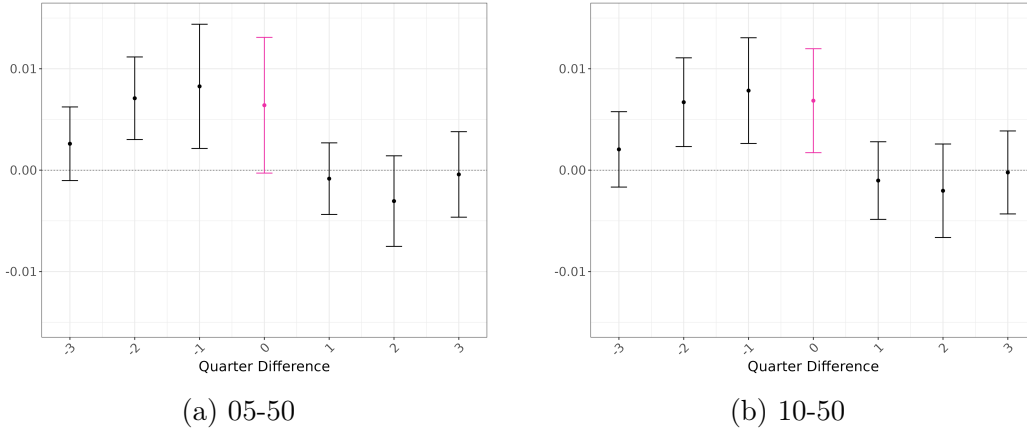


Figure L17: Event study estimates for π_{it} for countries using 05-50 bigrams threshold (left panel) and 10-50 bigrams threshold (right panel)

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from Specification 9 on the samples with different bigram selection thresholds. Left panel includes all bigrams used in at least 5 quarters with a maximum of at least 50 times per quarter. Right panel includes all bigrams used in at least 10 quarters with a maximum of at least 50 times per quarter. Errors are clustered at the elections level, 95% confidence intervals are reported. Outcome variable is populist partisanship, defined as in Equation 5.

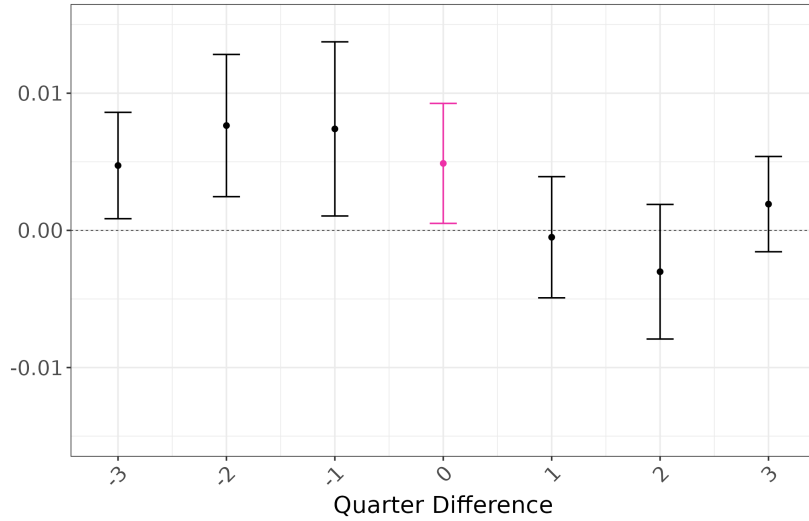


Figure L18: Event study estimates for π_{it} with ambiguous politicians coded as non-populists

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from Specification 9. Outcome variable is populist partisanship, defined as in Equation 5. Partisanship is estimated by assigning ambiguous politicians to the non-populist group (rather than to the populist group, as in the main text). Errors are clustered at the elections level, 95% confidence intervals are reported.

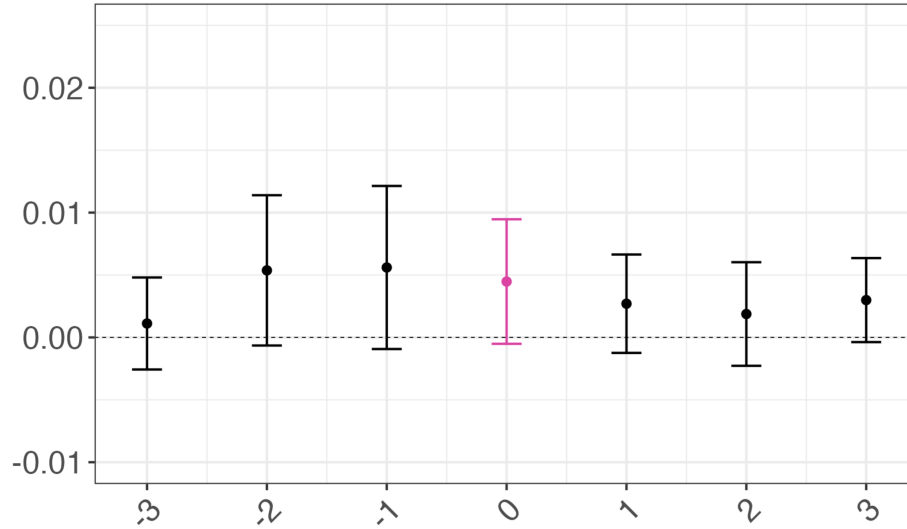


Figure L19: Event studies estimated on the set of non-classified bigrams

Notes. Point estimates and confidence intervals for $\beta_{d(i,t)}$ are plotted from Specification 9 on the sample of non-classified bigrams only (around 40% of all bigrams). Outcome variable is populist partisanship, defined as in Equation 5. Errors are clustered at the elections level, 95% confidence intervals are reported.

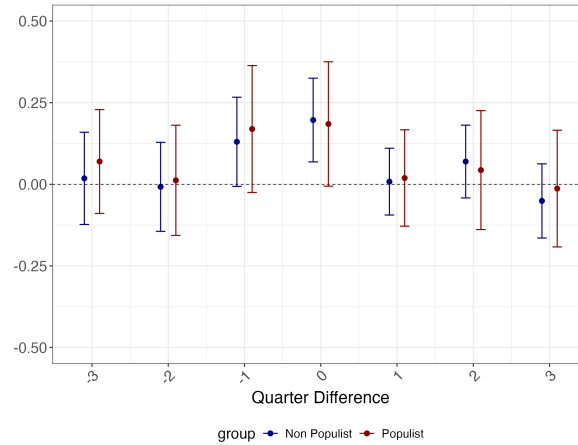


Figure L20: Event study estimates for the distinctiveness of the speech separately for populist and non-populist politicians

Notes. Figure plots point estimates and confidence intervals for $\beta_{d(i,t)}$ from specification 9 separately for populist and non-populist politicians. The outcome variable y_{it} is standardized and defined as in Equation 10. Errors are clustered at the elections level, 95% confidence intervals are reported.

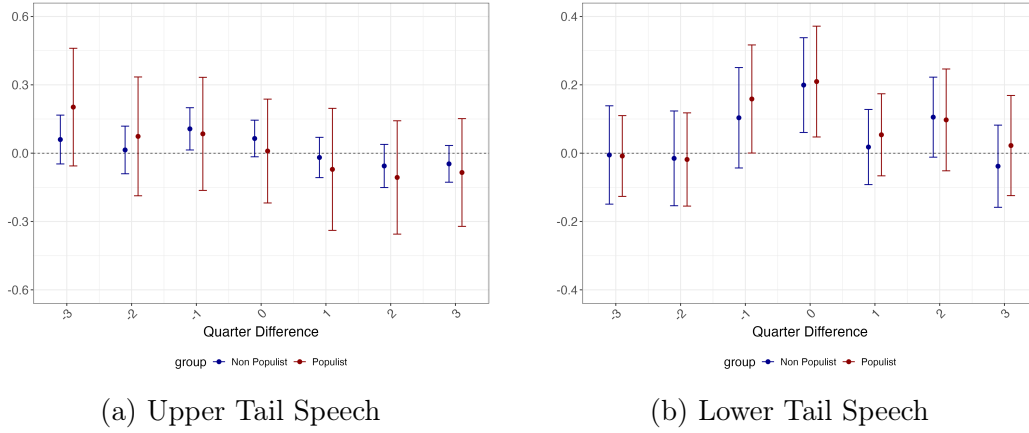


Figure L21: Event study estimates for the distinctiveness of tails of speeches separately for populist and non-populist politicians – Only Upper Tail Speech (Panel (a)), Only Lower Tail Speech (Panel (b)).

Notes. Point estimates and confidence intervals for $\beta_{d(i,t)}$ are plotted from specification 9 separately for populist and non-populist politicians on two subsamples: the left panel uses the most populist bigrams (with ξ_j above the 80th percentile of the ξ_j distribution), and the right panel uses the most non-populist bigrams (with ξ_j below the 20th percentile). The outcome variable y_{it} is standardized and defined as in Equation 10. Errors are clustered at the elections level, 95% confidence intervals are reported.

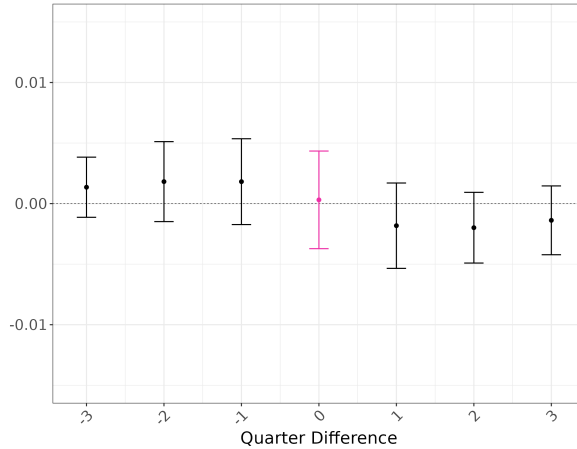


Figure L22: Event studies for π_{it} , Left-Right partisanship measure. Errors are clustered at the elections level, 95% confidence intervals are reported.

Notes. Point estimates and confidence intervals are plotted for $\beta_{d(i,t)}$ from Specification 9. Outcome variable is left-right partisanship bases with affiliation defined based on Manifesto RILE variable (instead of ChatGPT classification). Partisanship is defined as in Equation 5. Errors are clustered at the elections level, 95% confidence intervals are reported.