

Do politicians put their policies where their mouths are?

Germain Gauthier, Caroline Le Pennec, Vincent Pons

APSA Conference

September 2026

Motivation

“Read my lips: no new taxes.”

– George H. W. Bush, accepting the Republican nomination, 1988

1990: signs the Omnibus Budget Reconciliation Act – raising income, payroll, and excise taxes.

Are candidates' campaign statements a reliable guide to what they would do once in office?

Motivation

- ▶ Electoral campaign positions are strategic (Le Pennec, 2024; Di Tella et al., 2025) and, once the election is over, not legally binding.
- ▶ So, how informative and committal are candidate platforms? We consider three hypotheses:
 1. **Cheap talk.** Campaign positions contain no information about how a candidate would govern.
 2. **Signaling without commitment.** Platforms contain information about candidates' underlying policy preferences, but campaign-induced deviations from those preferences disappear after the election.
 3. **Commitment.** Campaign positions shape subsequent policy: even positions adopted strategically persist in office.
- ▶ Two empirical questions:
 - ▶ **Informativeness:** Do campaign positions predict in-office behavior? (2–3 vs. 1)
 - ▶ **Commitment:** Do campaign-induced changes in positions cause changes in in-office behavior? (3 vs. 2)

This talk

- ▶ Today: we focus on **informativeness**. (Commitment – the causal effect of the campaign on subsequent behavior – is work in progress.)
- ▶ Roadmap:
 1. Data: campaign and office text from France and the United States.
 2. Two text measures: ideological slant and 56 policy stances (e.g., supporting a free-market economy, defending labor groups, opposing multiculturalism, etc).
 3. The raw picture: we study how these measures evolve across the election cycle, from campaign to office.
 4. Prediction: we determine whether a politician's discourse during the campaign predicts their discourse once in office, both overall and within party.
 5. Causality: we use a close election RDD to test whether office discourse aligns with the platform of the winner better than the platform of the runner-up.

Outline

Data and Measures

Descriptive Evidence

Winner vs. Runner-Up RDD

Conclusion

Data: campaign text and office text, two countries

France – legislative elections, 1958–2022.

- ▶ *Campaign*: "professions de foi" (official two-pages manifestos issued by individual candidates), round 1 and round 2; 8006 races, 71059 candidacies.
 - ▶ Source: Di Tella et al. (2025)
- ▶ *Office*: the winner's parliamentary speeches over the same period.
 - ▶ Collected from digitized official records.

United States – House general elections, 2002–2018.

- ▶ *Campaign*: captures of archived campaign websites, primary and general stage; 3911 general races, 10333 general-election candidacies.
 - ▶ Source: Di Tella et al. (2025)
- ▶ *Office*: the winner's speeches in the Congressional Record.
 - ▶ Collected from digitized official records.

Measure 1: Ideological Slant (Gentzkow et al., 2019)

A common left-right scale across campaign and in-office communication

Pre-processing. We use an LLM to remove every mention of a candidate's name, party, or place, and curate the vocabulary to focus on unigrams and “real” bigrams.

Data. A document is defined as the aggregated text for a politician in a given election year and a given election stage.

- ▶ US: all website captures before the primary, all website captures between the primary and the general, and all interventions in Congress over the following term.
- ▶ France: round-1 manifesto, round-2 manifesto, and all interventions in parliament over the following term.

Each document d is represented as its count c_{dv} of uni-/bigrams $v \in \mathcal{V}$; $n_d = \sum_v c_{dv}$ is the length of the document.

Measure 1: Ideological Slant (Gentzkow et al., 2019)

A common left-right scale across campaign and in-office communication

Model. $c_{dv} \sim \text{Poisson}(\mu_{dv})$, with:

$$\log \mu_{dv} = \underbrace{\log n_d}_{\text{offset}} + \underbrace{\alpha_v}_{\text{phrase}} + \underbrace{\xi_{m(d),v}}_{\text{election stage}} + \underbrace{\lambda_{t(d),v}}_{\text{election year}} + \underbrace{g_d \phi_v}_{\text{ideology}} \quad (\text{M1})$$

$g_d = \mathbb{1}\{\text{right bloc}\}$ is the politician's party side, so ϕ_v is the right-minus-left loading of phrase v . Estimation uses all documents whose politician has a clear left/right label: Republican vs. Democrat in the US, right-wing vs. left-wing in France (centrists, third parties, and unclassified candidates are excluded from the estimation but scored afterward). ξ and λ absorb election stage- and election year-specific language during estimation, and no covariate enters at scoring time $\Rightarrow$ campaign and office documents are scored on one scale.

Main measure. Every document is scored by projecting its phrase shares onto the loadings:

$$\text{Ideology slant}_d = \sum_{v \in \mathcal{V}} \frac{c_{dv}}{n_d} \phi_v \quad (\text{M1}')$$

Positive = more right-coded language. We refer to it as the document's *slant*.

Ideological Phrases – France

The most partisan phrases in the estimated vocabulary



Notes: Left panel: most left-leaning phrases ($\hat{\phi}_v < 0$); right panel: most right-leaning ($\hat{\phi}_v > 0$). Phrases must occur in at least 50 documents in both campaign and office corpora; selection and size use $|\hat{\phi}_v| \log(1 + df_v)$, where df_v is the smaller of the two source-specific document frequencies.

Measure 2: Policy Stances

Every sentence classified into one of the 56 Manifesto Project policy categories

Data. The raw text W_d of each document is split into sentences s .

Filtering. Many sentences do not provide any information on the politician's policy stances (e.g., thanking the voters for their support, greeting fellow MPs in parliament, etc). We manually annotated a sample of 3769 sentences across all corpora (campaign and office text) and fine-tuned a multilingual classification model (mmBERT) to retain only policy-related sentences $\mathcal{R}_d$ (held-out F1 = .81).

Classifier. We leverage the pre-trained Manifesto Project's multilingual sentence classifier (fine-tuned on ~ 1.6 m hand-annotated statements), which gives each policy-related sentence a probability distribution over $C = 56$ categories defined by the Manifesto Project (see the full list on the next slide), and a hard label:

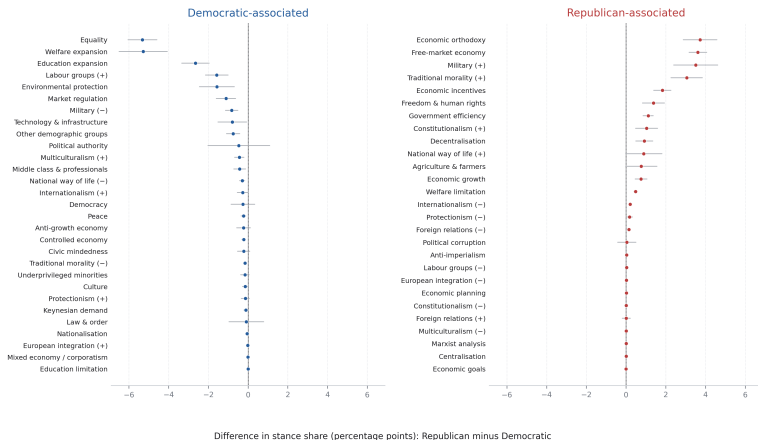
$$\pi_{dsc} = \Pr(c \mid W_{ds}), \quad c_{ds}^* = \arg \max_c \pi_{dsc} \quad (\text{M2})$$

Main measure. The share of document d 's policy-related sentences assigned to category c :

$$\hat{s}_{dc} = \frac{1}{|\mathcal{R}_d|} \sum_{s \in \mathcal{R}_d} \mathbb{1}\{c_{ds}^* = c\} \quad (\text{M2}')$$

Do the Stances Separate the Parties? – United States

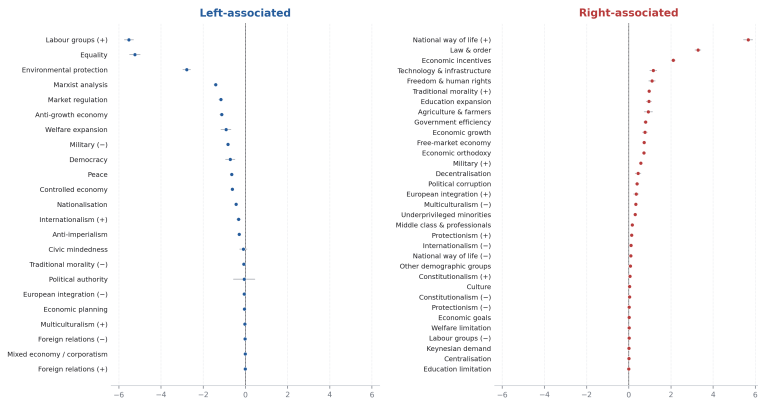
Estimated partisan gap in stance shares, by category (percentage points; Republican – Democrat)



Notes: For each category c , we estimate $\hat{s}_{dc} = \theta_c g_d + \xi_{m(d)} + \lambda_{t(d)} + \varepsilon_{dc}$ on the sample of politicians who are either Republican or Democrat and plot the party coefficients $100\hat{\theta}_c$. Regressions include election-stage and election-year fixed effects. Sample: 5070 documents (campaign websites and Congress); 95% CIs clustered by politician (833).

Do the Stances Separate the Parties? – France

Estimated partisan gap in stance shares, by category (percentage points; right – left)



Difference in stance share (percentage points): Right minus Left

Notes: For each category c , we estimate $\hat{s}_{dc} = \theta_c g_d + \xi_{m(d)} + \lambda_{t(d)} + \varepsilon_{dc}$ on the sample of politicians with a clear left/right label and plot the coefficients $100\hat{\theta}_c$. Regressions include election-stage and election-year fixed effects. Sample: 43733 documents (manifestos and Parliament); 95% CIs clustered by politician (32200).

Outline

Data and Measures

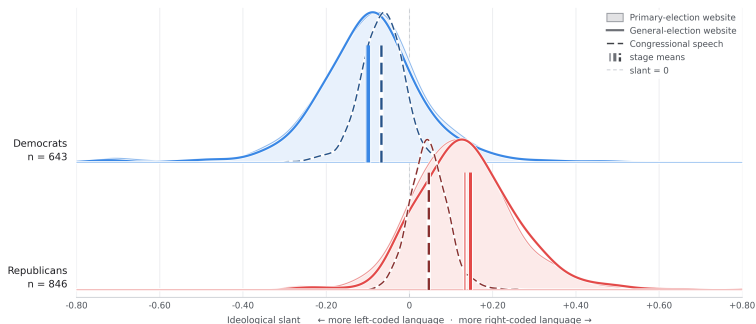
Descriptive Evidence

Winner vs. Runner-Up RDD

Conclusion

Slant across the Election Cycle – United States

The same winners at three stages: primary website, general-election website, congressional speech

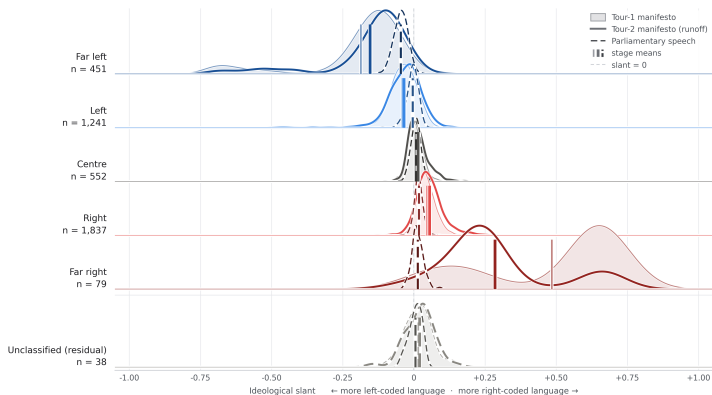


Notes: Kernel densities of the ideology score ($M1'$), by party and election stage, on the balanced sample of 1489 members of the House (643 Democrats, 846 Republicans) with a primary website, a general-election website, and speeches in the following Congress; cycles 2002–2016. Vertical bars mark stage means.

The moderation once in office is not a statistical artifact.

Slant across the Election Cycle – France

The same winners at three stages: round-1 manifesto, round-2 manifesto, parliamentary speech



Notes: Kernel densities of the ideology score (M1'), by election stage, on the balanced sample of 4198 members of parliament with a round-1 manifesto, a round-2 manifesto, and speeches in the following parliamentary term; elections 1958–2022. The six rows are the party-family categories, from far left to far right; the bottom row is nonclassified ($n = 38$). Vertical bars mark stage means.

The moderation once in office is not a statistical artifact.

Do Campaign Positions Predict Office Positions?

- ▶ We find evidence of discourse moderation from first to second round (Le Pennec (2024); Di Tella et al. (2025)), and even more so from second round to office → office discourse is different from campaign discourse.
- ▶ But we also find that politicians tend to remain on the same ideological side → campaign and office speech may still be correlated.
- ▶ We test this hypothesis by estimating, for each campaign stage separately,

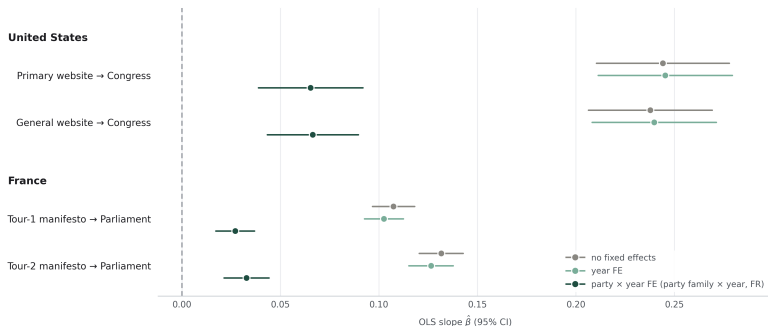
$$\text{Office}_{ie} = \beta \text{Campaign}_{ie} + \varepsilon_{ie}$$

where Office_{ie} is the slant (M1') of politician i 's speeches in the parliamentary term following election year e , and Campaign_{ie} is the slant of that same politician's campaign text in election year e .

- ▶ We add election year fixed effects and party $\times$ year fixed effects to test whether campaign discourse predicts office speech beyond what global politics (e.g., all political discourse shifts to the right over time) or parties (e.g., Republican politicians always use more right-wing language) would predict.

Do Campaign Positions Predict Office Positions? Slant

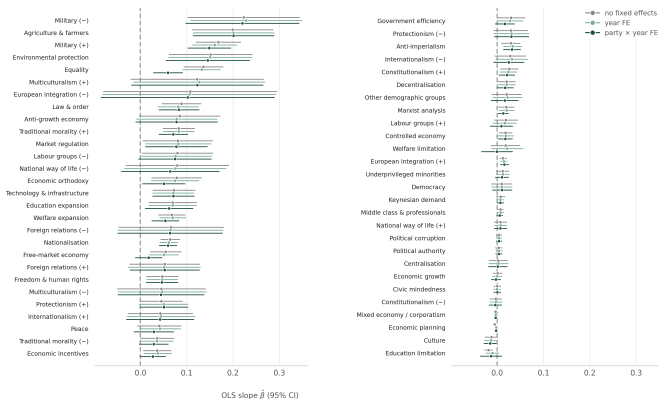
Estimated slope of in-office slant on campaign slant, with and without party $\times$ year fixed effects



Notes: Specifications: no FE, election-year FE, and party $\times$ year FE (party family $\times$ year in France). Sample sizes (party-FE sample in parentheses): US primary 1793 (1791) cells, general 1720 (1718) / 755 (754); France round 1 5119 (5115) / 3160 (3158), round 2 4370 (4366) / 2969 (2968). 95% CIs clustered by politician. Within party $\times$ year, $\hat{\beta} = .065-.066$ (US) and $.027-.033$ (France), 72–75% below the raw slope; all $p < 10^{-5}$.

Do Campaign Stances Predict Office Stances? – US

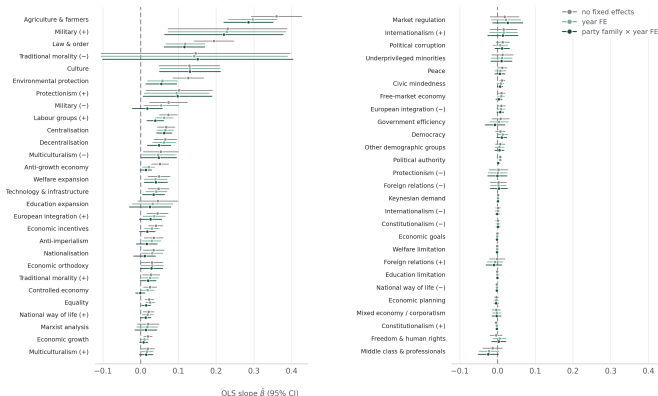
Estimated slope of office emphasis on campaign emphasis, by policy category (primary websites → Congress)



Notes: For each category c we estimate $\hat{s}_{iec}^{\text{office}} = \beta_c \hat{s}_{iec}^{\text{campaign}} + \text{FE} + \varepsilon_{iec}$ under the same three specifications as for ideological slant. Shares lie in $[0, 1]$, so β_c is a pass-through rate ($\beta_c = 1$: full pass-through). Categories are sorted by the no-FE slope; 23 of 55 retain a 95% CI excluding zero under party $\times$ year FE. Sample sizes: 1449 observations (1448 under party $\times$ year FE); CIs clustered by politician.

Do Campaign Stances Predict Office Stances? – France

Estimated slope of office emphasis on campaign emphasis, by policy category (round-1 manifestos → Parliament)



Notes: For each category c we estimate $\hat{s}_{iec}^{\text{office}} = \beta_c \hat{s}_{iec}^{\text{campaign}} + \text{FE} + \varepsilon_{iec}$ under the same three specifications as for ideological slant. Shares lie in $[0, 1]$, so β_c is a pass-through rate ($\beta_c = 1$: full pass-through). Categories are sorted by the no-FE slope; 15 of 55 retain a 95% CI excluding zero under party-family $\times$ year FE. Sample: 4998 cells / 3118 politicians (4995 / 3116 under party-family $\times$ year FE); CIs clustered by politician.

What do we learn from these correlations?

- ▶ **Party explains a lot, but not everything.** Party $\times$ year fixed effects absorb 72–75% of the campaign–office relationship in ideological slant. For policy stances, substantial within-party variation remains.
→ Candidate platforms provide more information than simply knowing the politician's party.
- ▶ **This could reflect candidate preferences, but also district heterogeneity.** Candidates from the same party run in different districts: district characteristics may shape both the campaign positions of *all* candidates (e.g., candidates in rural districts talk more about agriculture) and the elected politician's behavior in office. Observing candidate platforms may not add more information than knowing their district.
- ▶ **Thus, we rely on within-race variation.** In close elections, we can compare the winner's subsequent office text with the campaign platforms of the winner and the runner-up, whilst holding the electoral environment fixed (and how well the candidates cater to their constituents). Is the narrow winner's platform more informative than the platform of his narrow runner-up?

Outline

Data and Measures

Descriptive Evidence

Winner vs. Runner-Up RDD

Conclusion

Winner vs. Runner-Up RDD

Design. We use two observations per race, corresponding to the winner and runner-up in the final round. We compare the campaign platforms of both with the winner's subsequent office text.

Outcome. A similarity S between the winner's office text o_e and candidate i 's campaign text p_i :

$$Y_i = S(o_e, p_i), \quad \begin{cases} S^{\text{slant}}(o_e, p_i) = -|o_e - p_i|, \\ S^{\text{stances}}(o_e, p_i) = \frac{\tilde{o}_e' \tilde{p}_i}{\|\tilde{o}_e\| \|\tilde{p}_i\|}. \end{cases}$$

where $\tilde{x}$ is x minus its election-year $\times$ election-stage mean.

→ Higher values mean greater campaign–office proximity.

Estimation. With X_i the difference in vote shares between the winner and the runner-up (positive for the winner, negative for the runner-up), and $T_i = \mathbb{1}\{i \text{ won}\}$:

$$Y_i = \alpha + \tau T_i + \beta X_i + \gamma T_i X_i + \varepsilon_i. \quad (\text{RDD})$$

We use a local linear, triangular kernel; robust bias-corrected 95% CIs; SEs clustered by district $\times$ year. The sample is restricted to pairs where we could score platforms for both finalists, and a party benchmark built from other candidates of the same party and election year is available for each.

sample definition

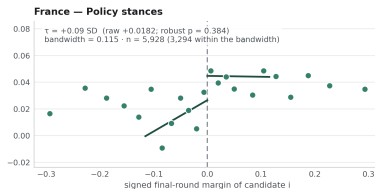
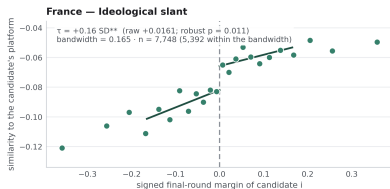
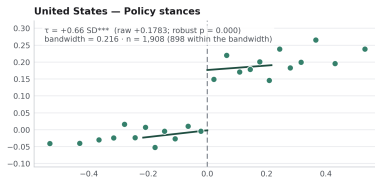
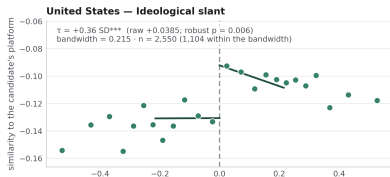
diagnostics

all races

naive OLS

The Discontinuities in the Raw Data

Binned means of campaign–office similarity against the winner's margin, by country and measure



Notes: Dots are binned means of $Y_i = S(o_e, p_i)$ against candidate i 's vote share difference with the runner-up (for winners) or with the winner (for runner-ups); lines are local-linear fits within the MSE-optimal bandwidth (triangular kernel). $\hat{\tau}$ is the jump at the cutoff in SDs of the outcome; p is robust bias-corrected. Sample: main RDD sample with two observations per race. n gives the number of plotted observations, with observations inside the bandwidth in parentheses.

A Decomposition Exercise

The platform of the winner is more informative than the platform of an equally-good candidate. **Is this winner effect only about the winner's party, or also about the candidate?**

Let J_i be the set of other final-round candidates from i 's party in the same election year. Then we can write:

$$\underbrace{S(o_e, p_i)}_{Y_i^{\text{total}}} = \underbrace{\frac{1}{|J_i|} \sum_{j \in J_i} S(o_e, p_j)}_{Y_i^{\text{party}}} + \underbrace{\left[S(o_e, p_i) - \frac{1}{|J_i|} \sum_{j \in J_i} S(o_e, p_j) \right]}_{Y_i^{\text{cand}}}$$

We fix the kernel and the bandwidth of an RDD on Y_i^{total} , then run the same specification on Y_i^{party} and Y_i^{cand} . We obtain:

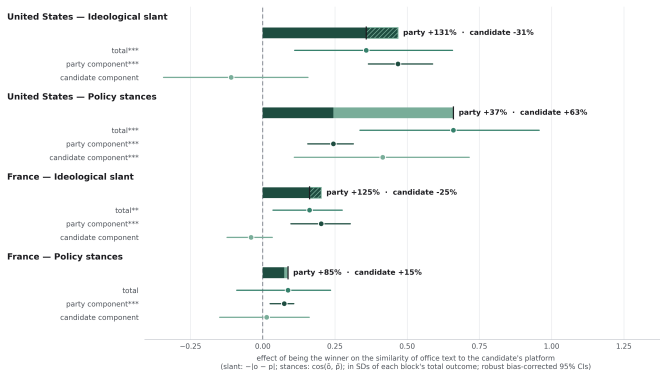
$$\hat{\tau}_{\text{total}} = \hat{\tau}_{\text{party}} + \hat{\tau}_{\text{cand}}.$$

In other words, we can disentangle whether the candidate's platform is informative because office speech aligns with that candidate's specific campaign discourse, or because it aligns with the candidate's overall party discourse.

decomposition details

Decomposition Results

The winner effect split into a party and a candidate component, by country and measure



Notes: Three local-linear RDs use the same main RDD sample observations, triangular kernel, and the total outcome's MSE-optimal bandwidth, so the split is exact: $\hat{\tau}_{\text{total}} = \hat{\tau}_{\text{party}} + \hat{\tau}_{\text{candidate}}$. Effects are in SDs of each block's total outcome (slant: $-|o - p|$; stances: $\cos(\delta, \tilde{p})$). Negative components are hatched over the bar segment they offset; the tick marks the total. Robust bias-corrected 95% CIs are centered on the bias-corrected estimate and can therefore be asymmetric around the conventional point. Shares are reported for every block and are not restricted to $[0, 1]$. Sample: slant 3874 France / 1275 US races; stances 2964 / 954. Stars: * $p < .1$, ** $p < .05$, *** $p < .01$ (robust p).

Outline

Data and Measures

Descriptive Evidence

Winner vs. Runner-Up RDD

Conclusion

Concluding Remarks

- ▶ **Campaign platforms are more than cheap talk.** Ideological slant and policy emphasis predict subsequent office discourse in both France and the United States.
- ▶ **Much of this information is about parties.** For ideological slant, party $\times$ year absorbs roughly three-quarters of the raw campaign–office relationship.
- ▶ Close elections let us isolate what is specific to the candidate's preferences. **Candidates' platforms are informative beyond the district characteristics, but mostly because office discourse aligns with the average discourse of the candidate's party.**
- ▶ One exception: **in the United States, policy stances strongly reflect candidates' preferences beyond their party affiliation.**
- ▶ Next steps: exploit variation in the identity of the winner's opponent during the campaign and the induced strategic adjustments of winners' platforms to test whether they commit to these adjustments.

FR Manifesto Coverage by Election

Share of candidacies with a scored round-1 profession de foi (all candidacies / elected MPs)

Election	all	elected	Election	all	elected
1958	67%	70%	1993	90%	94%
1962	76%	48%	1997	85%	92%
1967	90%	92%	2002	2%	6%
1968	95%	96%	2007	5%	8%
1973	91%	95%	2012	5%	5%
1978	90%	93%	2017	63%	72%
1981	89%	95%	2022	75%	86%
1988	82%	84%			

Notes: Each cell is the share with a scored round-1 manifesto among all candidacies and among elected MPs in that election. Universe: 71059 legislative candidacies, 1958–2022; overall coverage is 56% for all candidacies and 68% for elected MPs. Manifestos for the 2002–2012 period are largely absent from the digitized corpus due to a change in data collection policy at the CEVIPOF (see Le Pennec (2024) for more details).

US Website Coverage by Cycle

Contested House seats by number of top-two general-election websites scored

Cycle	seats	both	one	neither	any (%)
2002	398	102	152	144	64
2004	400	152	177	71	82
2006	394	162	167	65	84
2008	403	151	185	67	83
2010	421	205	177	39	91
2012	417	200	167	50	88
2014	400	172	183	45	89
2016	401	169	184	48	88
2018	416	7	29	380	9

Notes: Universe: the 3650 contested House seats whose winner's congressional speech is scored, by cycle. "Both," "one," and "neither" count scored websites among the top two general-election candidates; "any" is the share with at least one. Congressional speech is scored for 3650 of 3666 matched winners (99.6%). The 2018 website gap reflects the archive itself; the RDD coverage diagnostic shows scoring does not jump at the cutoff.

Office Moderation Is Real: The Compression Survives Every Estimation

Campaign-to-office slant dispersion under alternative estimations of the ideology score

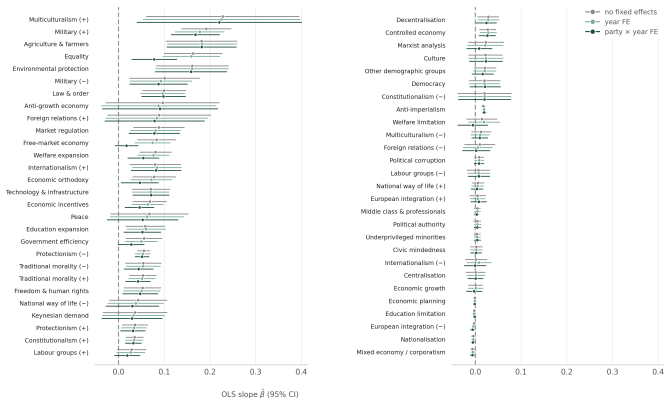
sd(campaign)/ sd(office) under. . .	France	United States
Pooled estimation (the deck's measure)	5.7	2.1
Token-mass-balanced estimation	6.5	2.5
Campaign-only estimation	5.8	3.0
Office-only estimation	2.1	1.7
Stance-based score (56 issue shares, no phrases)	2.2	2.2

- ▶ **Vocabulary overlap is not the issue:** $\geq 99.4\%$ of each source's token mass and $\geq 96.6\%$ of its $|\phi|$ -mass fall on phrases shared with the other source (both countries).
- ▶ **Rankings are also stable across phrase-based fits:** rank correlations of document scores are 0.80–0.99 between the balanced fit and any other phrase-based fit.

Even a model estimated only on office speech spreads campaign text 1.7–2.1 $\times$ more, and a score that sees only the 56 issue shares – no phrases at all – still compresses office speech.

Do Campaign Stances Predict Office Stances? – US

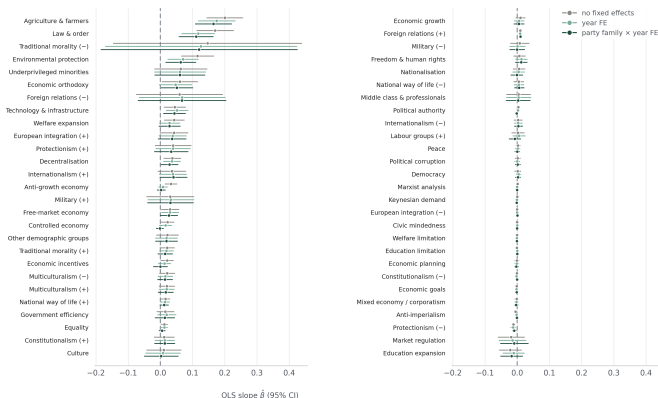
Estimated slope of office emphasis on campaign emphasis, by policy category (general websites → Congress)



Notes: One OLS regression per category under no FE, election-year FE, and party $\times$ year FE; shares lie in $[0, 1]$, so the coefficient is a pass-through rate. Categories are sorted by the no-FE slope; 26 of 55 retain a 95% CI excluding zero under party $\times$ year FE. Sample: 1423 cells / 688 politicians (1422 / 687 under party $\times$ year FE); CIs clustered by politician.

Do Campaign Stances Predict Office Stances? – France

Estimated slope of office emphasis on campaign emphasis, by policy category (round-2 manifestos → Parliament)



Notes: One OLS regression per category under no FE, election-year FE, and party-family × year FE; shares lie in $[0, 1]$, so the coefficient is a pass-through rate. Categories are sorted by the no-FE slope; 12 of 54 retain a 95% CI excluding zero under party-family × year FE. Sample: 3705 cells / 2701 politicians (3700 / 2698 under party-family × year FE); CIs clustered by politician.

The RDD Sample Funnel

From all races to the main RDD sample

France	counts: races
All legislative races, 1958–2022	8,006
decided in a runoff, both finalists' round-2 vote shares	6,811
winner's parliamentary speech scored	6,489
both-finalists sample: both round-2 manifesto slants scored	3,946
both-finalists sample: both round-2 stance vectors available	3,018
main RDD sample: party benchmark available for both finalists (slant / stances)	3,874 / 2,964
United States	counts: seats
All House general elections, 2002–2018	3911
contested (walkovers have no runner-up: 239)	3,672
winner matched to the serving member	3,666
winner's congressional speech scored	3,650
both-finalists sample: both top-two website slants scored	1,320
both-finalists sample: both top-two website stance vectors available	982
main RDD sample: party benchmark available for both finalists (slant / stances)	1,275 / 954

Notes: Each retained race contributes two candidate observations. A stance vector exists only for platforms with informative sentences. The main RDD sample also requires a leave-one-out party benchmark for each finalist, built from other scored platforms in the same party and election year (party family in France); a race drops if either finalist lacks the required party class.

The Mirror Regression Without the Threshold

OLS winner advantage on all races and the main RDD sample vs. the RD at the cutoff

Losers may lose precisely because their platform fit the district less – the naive winner advantage can therefore overstate informativeness (and indeed OLS estimates are generally larger).

		OLS, all races	OLS, main RDD sample	RD at cutoff	races (all / main)
FR	slant	+0.39***	+0.38***	+0.16**	4,608 / 3,874
FR	stances	+0.13***	+0.11***	+0.09	4,188 / 2,964
US	slant	+0.27***	+0.28***	+0.36***	2,682 / 1,275
US	stances	+0.76***	+0.80***	+0.66***	2,419 / 954

Notes: The winner's office text is paired with the winner's and the runner-up's platform. OLS regresses Y on the winner indicator without a bandwidth; "all races" admits races where we miss one of the two finalists' platform, whereas the main RDD sample requires both finalists' platforms and a party benchmark for each. Outcomes are $-|o - p|$ and $\cos(\vec{o}, \vec{p})$, in SDs of the outcome. OLS uses district $\times$ year cluster-robust SEs on the corresponding RD observations; RD entries use robust bias-corrected p -values on the main RDD sample.

RDD Diagnostics

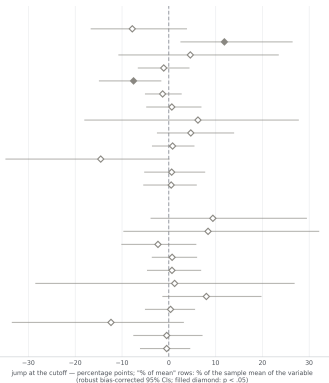
RD estimates on predetermined variables: all races vs. the main RDD sample

All races

US — candidate's platform is scored
 $p = 0.22$
US — candidate is a Democrat
 $p = 0.02$
US — platform-deviation norm (% of mean)
 $p = 0.46$
FR — candidate's manifesto is scored
 $p = 0.69$
FR — candidate is left-wing (left vs right bloc)
 $p = 0.01$
FR — candidate is female
 $p = 0.54$
FR — platform-deviation norm (% of mean)
 $p = 0.72$
US — [p - q], candidate to party benchmark (% of mean)
 $p = 0.68$
US — cos(β , q), candidate's direction vs party benchmark (x 100)
 $p = 0.17$
US — other candidates in party $\times$ year, K (% of mean)
 $p = 0.68$
FR — [p - q], candidate to party benchmark (% of mean)
 $p = 0.25$
FR — cos(β , q), candidate's direction vs party benchmark (x 100)
 $p = 0.79$
FR — other candidates in party $\times$ year, K (% of mean)
 $p = 0.93$

Main RDD sample

US — candidate is a Democrat
 $p = 0.13$
US — platform-deviation norm (% of mean)
 $p = 0.29$
FR — candidate is left-wing (left vs right bloc)
 $p = 0.60$
FR — candidate is female
 $p = 0.62$
FR — platform-deviation norm (% of mean)
 $p = 0.70$
US — [p - q], candidate to party benchmark (% of mean)
 $p = 0.95$
US — cos(β , q), candidate's direction vs party benchmark (x 100)
 $p = 0.08$
US — other candidates in party $\times$ year, K (% of mean)
 $p = 0.93$
FR — [p - q], candidate to party benchmark (% of mean)
 $p = 0.10$
FR — cos(β , q), candidate's direction vs party benchmark (x 100)
 $p = 0.96$
FR — other candidates in party $\times$ year, K (% of mean)
 $p = 0.77$



Notes: Each row is a separate local-linear RD of a predetermined outcome, estimated at its own MSE-optimal bandwidth with a triangular kernel and robust bias-corrected 95% CI. Dummies are in percentage points; “% of mean” rows are relative to that sample’s outcome mean; a filled diamond marks robust $p < .05$. For the main RDD sample, we additionally test the candidate’s distance or direction relative to the party benchmark and the number K of other candidates in the same party and election year.

Decomposition Details

Conditional on a fixed common bandwidth and design matrix, the estimator is linear in the outcome.

With the (RDD) notation, kernel weights $w_i = K(X_i/h)$, regressors Z_i stacked into Z , $W = \text{diag}(w_1, \dots, w_n)$, and $e_2 = (0, 1, 0, 0)'$ selecting the coefficient on T_i , the local-linear estimate is weighted least squares – a fixed set of weights applied to the outcome:

$$\hat{\tau} = e_2'(Z'WZ)^{-1}Z'WY = \sum_i a_i Y_i, \quad a_i = e_2'(Z'WZ)^{-1}Z_i w_i.$$

The a_i depend on the margins, the kernel, and the bandwidth – never on the outcome. Same races and same bandwidth across the three regressions therefore means the *same* a_i applied to three different outcomes; and because $Y_{\text{cand},i} = Y_{\text{total},i} - Y_{\text{party},i}$ by construction, the outcomes add row by row,

$$\begin{aligned}\hat{\tau}_{\text{total}} &= \sum_i a_i (Y_{\text{party},i} + Y_{\text{cand},i}) \\ &= \sum_i a_i Y_{\text{party},i} + \sum_i a_i Y_{\text{cand},i} \\ &= \hat{\tau}_{\text{party}} + \hat{\tau}_{\text{cand}}.\end{aligned}$$

The Decomposition on the All-Races Samples

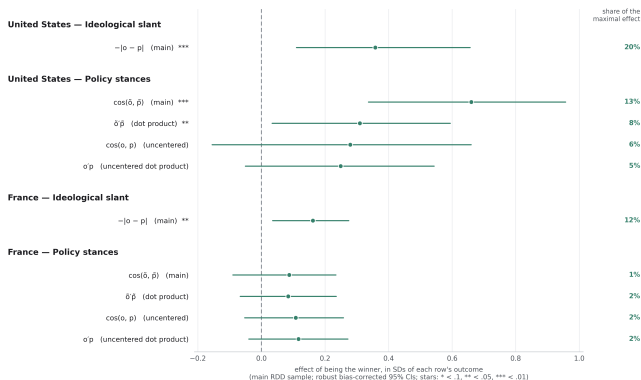
The decomposition without requiring both finalists to be observed

		total $\hat{\tau}$ (SD)	party share	candidate share	races (obs.)
US	slant	+0.42***	+111%***	-11%	2,682 (3,957)
US	stances	+0.60***	+38%***	+62%***	2,419 (3,373)
FR	slant	+0.18***	+124%***	-24%	4,608 (8,482)
FR	stances	+0.13	+70%***	+30%	4,188 (7,152)

Notes: A candidate (winner or runner-up) is included in the estimation whenever its own office text, platform, and party benchmark exist, regardless of platform availability for the opposing finalist. Totals are in SDs of the outcome; party and candidate entries are their shares of the total and are not restricted to $[0, 1]$. Stars use robust p -values: * $p < .1$, ** $p < .05$, *** $p < .01$.

Alternative Similarity Measures, Same Design

The total effect under alternative similarity functions, in SDs and as a share of the maximal effect



Notes: Main RDD sample. Each row is one similarity function, showing the total effect $\hat{\tau}$ in SDs of that outcome; robust bias-corrected 95% CIs. Right column: $\hat{\tau}$ as a share of the maximal effect attainable at the threshold. Slant row: the absolute distance (main; cosine is degenerate in one dimension). Stance rows: the cosine of the centered stance-share vectors (main), their dot product, and the cosine and dot product of the *uncentered* share vectors.

Why That Is the Maximal Effect

Each $\hat{\tau}$, benchmarked against the largest effect the design could deliver

The idea. At the cutoff the RD compares the *same* office text against the two platforms. However informative campaigns are, that comparison has a mechanical ceiling: how far apart the two platforms were. The printed share is $\hat{\tau}$ over that ceiling.

Distance (slant, uncentered). Put the two platforms on the line, $D = |p_w - p_l|$ apart, and let office speech land at o :

$$Y_w - Y_l = |o - p_l| - |o - p_w| \in [-D, +D].$$

- ▶ o at the winner's platform – or anywhere beyond it, away from the loser: the difference is exactly D . The full gap transmits, and overshooting cannot beat it.
- ▶ o between the platforms, δ from p_w : the difference is $D - 2\delta$ – positive while o is closer to the winner's platform, negative past the midpoint.

Race by race the effect is capped by the platform gap, so τ is capped by the gap's level at the cutoff (winner-side local-linear limit, same bandwidth as $\hat{\tau}$).

Same logic, other geometries. Cosine: the best any office *direction* can do is $\|\hat{p}_w - \hat{p}_l\|$. Dot product: $Y_w - Y_l = o'(p_w - p_l)$, so given the office magnitude the ceiling is $\|o\| \cdot \|p_w - p_l\|$.

Robustness: Replace the Winner's Office Text

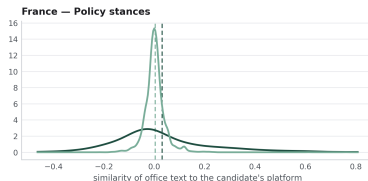
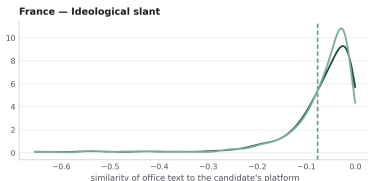
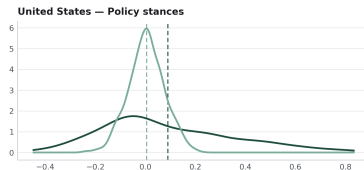
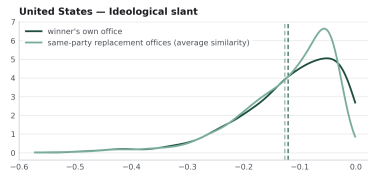
Does the winner's own office text track the platform better than a co-partisan's?

		own $\hat{\tau}$	replacement $\hat{\tau}$	difference (robust p)	races (obs.)
US	slant	+0.36***	+0.50***	-0.14 (.006)	1,275 (2,550)
US	stances	+0.66***	+0.35***	+0.31 (.053)	954 (1,908)
FR	slant	+0.16**	+0.16**	-0.00 (.896)	3,872 (7,744)
FR	stances	+0.09	+0.09***	-0.01 (.818)	2,962 (5,924)

Notes: We re-estimate each main-outcome RD after replacing the winner's office text by every other winner from the same party and election year with scored office text; each replacement is a single document, and its similarities are averaged only afterward. The difference is an RD on own minus replacement similarity. Main RDD sample plus at least one replacement office text; effects are in SDs of the own outcome. All three RDs use the same observations and the own-text fit's bandwidth, so the difference equals own minus replacement.

Robustness: Replace the Winner's Office Text

Similarity to the platform: the winner's own office text vs. same-party replacements



Notes: Main RDD sample, all observations. Dark: similarity of the winner's own office text to the row's platform (slant $-|o - p|$, stances $\cos(\vec{o}, \vec{p})$); light: the average, over the other winners from the same party and election year, of the same similarity computed with that winner's office text (dashed lines: means). Samples: US 1275 slant / 954 stance races; France 3874 / 2964. The replacement distribution is tighter because it averages similarities over many replacement texts.

References

- Di Tella, R., Kotti, R., Le Pennec, C., and Pons, V. (2025). Keep your enemies closer: Strategic platform adjustments during US and French elections. *American Economic Review*, 115(8):2488–2528.
- Gentzkow, M., Shapiro, J. M., and Taddy, M. (2019). Measuring group differences in high-dimensional choices: method and application to congressional speech. *Econometrica*, 87(4):1307–1340.
- Le Pennec, C. (2024). Strategic campaign communication: Evidence from 30,000 candidate manifestos. *The Economic Journal*, 134(658):785–810.