

Which Issues Divide?

Interpretable Ideal-Point Estimation from Political Text

Jefferson Leal

September 2, 2026

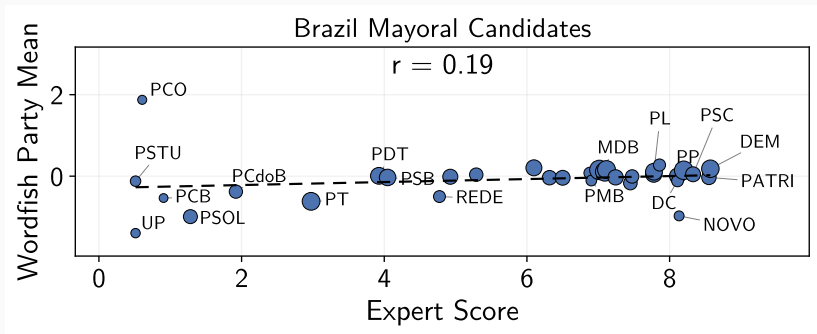
University of Rochester

When Scaling Ideology from Text Goes Wrong

- **Motivation:** Brazilian mayoral candidates campaign platforms 2020.

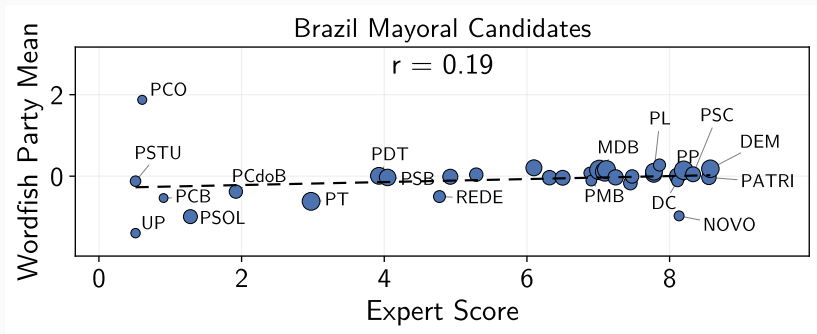
When Scaling Ideology from Text Goes Wrong

- Motivation:** Brazilian mayoral candidates campaign platforms 2020.



When Scaling Ideology from Text Goes Wrong

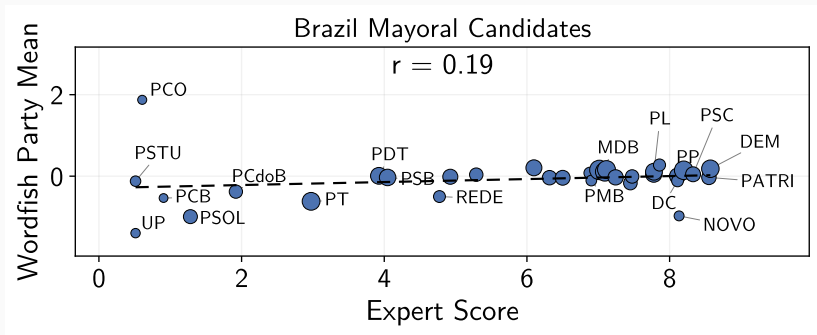
- **Motivation:** Brazilian mayoral candidates campaign platforms 2020.



- Supervised learning? divides must be known in advance.

When Scaling Ideology from Text Goes Wrong

- **Motivation:** Brazilian mayoral candidates campaign platforms 2020.



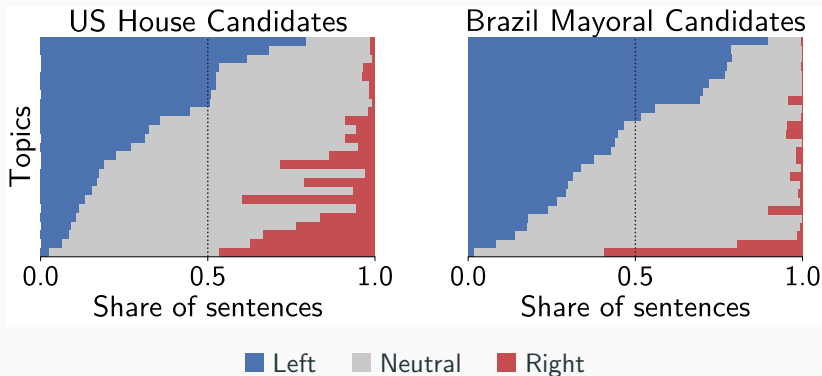
- Supervised learning? divides must be known in advance.
- Deep learning or LLMs? Black boxes.

Problem: Non-Polarized Language

- **Data:** 4,507 US House platforms, 2018–2022 primaries ([Porter et al., 2025](#)); 16,830 Brazilian mayoral platforms, 2020 ([Tribunal Superior Eleitoral, 2020](#)).

Problem: Non-Polarized Language

- **Data:** 4,507 US House platforms, 2018–2022 primaries ([Porter et al., 2025](#)); 16,830 Brazilian mayoral platforms, 2020 ([Tribunal Superior Eleitoral, 2020](#)).



Problem: Non-Polarized Language

- Global South: valence issues, i.e., public services & corruption control ([Kitschelt, 2007](#); [Bleck and van de Walle, 2012](#)).

Problem: Non-Polarized Language

- Global South: valence issues, i.e., public services & corruption control ([Kitschelt, 2007](#); [Bleck and van de Walle, 2012](#)).
- Local politics in the US is similar ([Bucchianeri, 2020](#)).

- Candidates reveal ideology through what issues they emphasize (**salience**) and the stated position per issue (**stance**) ([Laver, 2001](#); [Lauderdale and Herzog, 2016](#)).

Pipeline

- Candidates reveal ideology through what issues they emphasize (**salience**) and the stated position per issue (**stance**) ([Laver, 2001](#); [Lauderdale and Herzog, 2016](#)).
- **Pipeline:** sentences $\rightarrow$ topics + stances $\rightarrow$ joint IRT.

Pipeline

- Candidates reveal ideology through what issues they emphasize (**salience**) and the stated position per issue (**stance**) ([Laver, 2001](#); [Lauderdale and Herzog, 2016](#)).
- **Pipeline:** sentences $\rightarrow$ topics + stances $\rightarrow$ joint IRT.
- **Topics:** BERTopic ([Grootendorst, 2022](#)).

Pipeline

- Candidates reveal ideology through what issues they emphasize (**salience**) and the stated position per issue (**stance**) ([Laver, 2001](#); [Lauderdale and Herzog, 2016](#)).
- **Pipeline:** sentences $\rightarrow$ topics + stances $\rightarrow$ joint IRT.
- **Topics:** BERTopic ([Grootendorst, 2022](#)).
- **Stances:** NLI transformer ([Laurer et al., 2024](#)) fine-tuned with data from Manifesto Project ([Lehmann et al., 2025](#)).

Pipeline

- Candidates reveal ideology through what issues they emphasize (**salience**) and the stated position per issue (**stance**) ([Laver, 2001](#); [Lauderdale and Herzog, 2016](#)).
- **Pipeline:** sentences $\rightarrow$ topics + stances $\rightarrow$ joint IRT.
- **Topics:** BERTopic ([Grootendorst, 2022](#)).
- **Stances:** NLI transformer ([Laurer et al., 2024](#)) fine-tuned with data from Manifesto Project ([Lehmann et al., 2025](#)).
 - Keep only sentences where GPT-4.1 agrees with the human coder ([Burnham et al., 2026](#)), to mitigate low intercoder reliability ([Mikhaylov et al., 2012](#)).

Pipeline

- Candidates reveal ideology through what issues they emphasize (**salience**) and the stated position per issue (**stance**) ([Laver, 2001](#); [Lauderdale and Herzog, 2016](#)).
- **Pipeline:** sentences $\rightarrow$ topics + stances $\rightarrow$ joint IRT.
- **Topics:** BERTopic ([Grootendorst, 2022](#)).
- **Stances:** NLI transformer ([Laurer et al., 2024](#)) fine-tuned with data from Manifesto Project ([Lehmann et al., 2025](#)).
 - Keep only sentences where GPT-4.1 agrees with the human coder ([Burnham et al., 2026](#)), to mitigate low intercoder reliability ([Mikhaylov et al., 2012](#)).
 - Held-out balanced accuracy 0.83.

Fasti: From Attention, Stance and Topics to Ideal Points

- Per document and topic: a **sentence count** (salience) and a **mean stance**.

Fasti: From Attention, Stance and Topics to Ideal Points

- Per document and topic: a **sentence count** (salience) and a **mean stance**.
- One ideal point per candidate generates both.

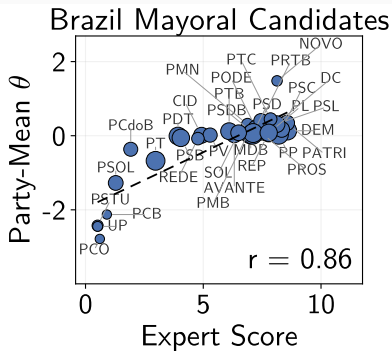
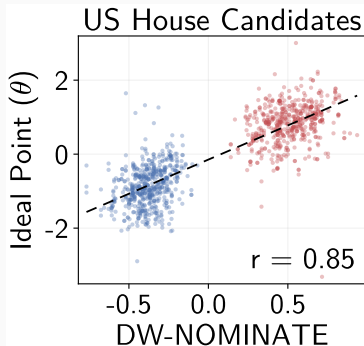
Fasti: From Attention, Stance and Topics to Ideal Points

- Per document and topic: a **sentence count** (salience) and a **mean stance**.
- One ideal point per candidate generates both.
- **Stance**: a linear latent factor model on ideal point.

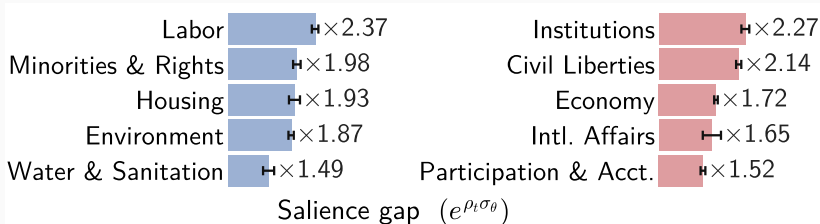
Fasti: From Attention, Stance and Topics to Ideal Points

- Per document and topic: a **sentence count** (salience) and a **mean stance**.
- One ideal point per candidate generates both.
- **Stance**: a linear latent factor model on ideal point.
- **Salience**: counts are multinomial across topics, each topic's expected share quadratic in the ideal point.

External Benchmarks

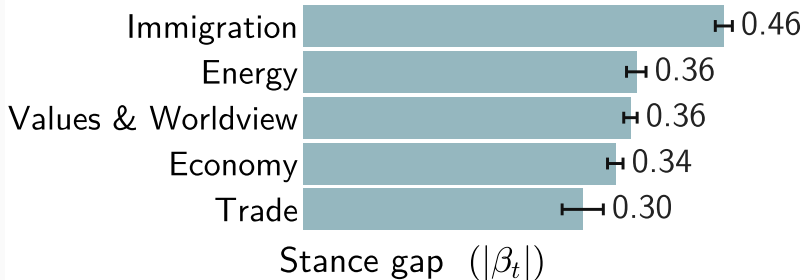


Salience Divisions: United States

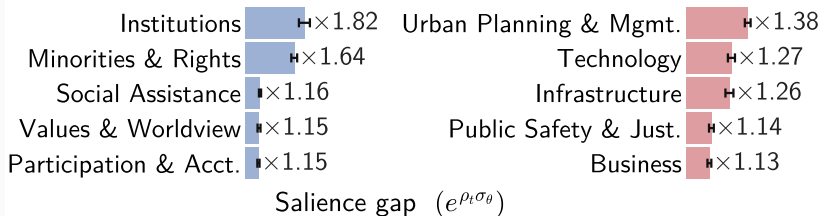


■ Left-emphasized ■ Right-emphasized

Stance Divisions: United States

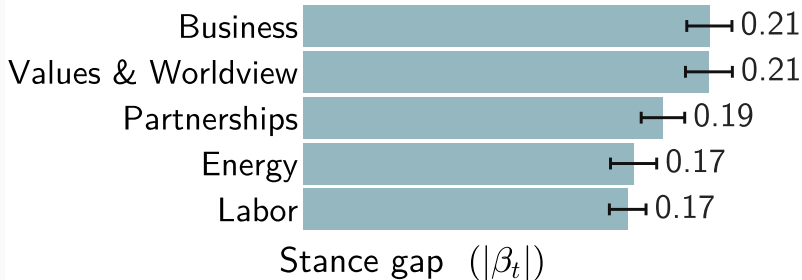


Salience Divisions: Brazil



■ Left-emphasized ■ Right-emphasized

Stance Divisions: Brazil



Candidate Strategic Positioning

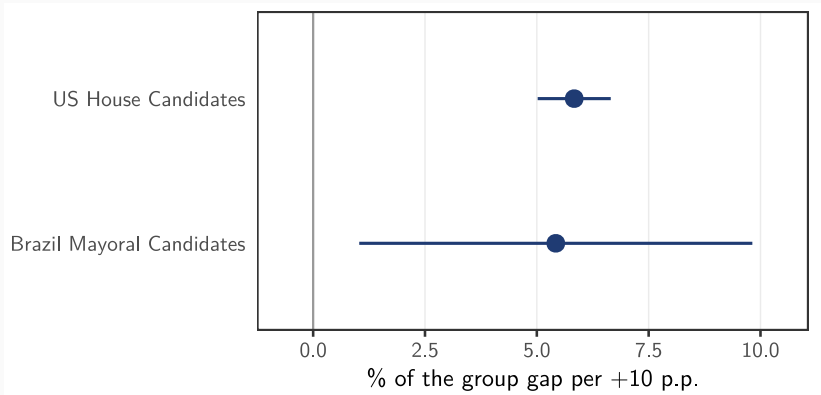
- Candidates strategically reposition to match the leaning of the district on presidential elections ([Meisels, 2026](#)).

Candidate Strategic Positioning

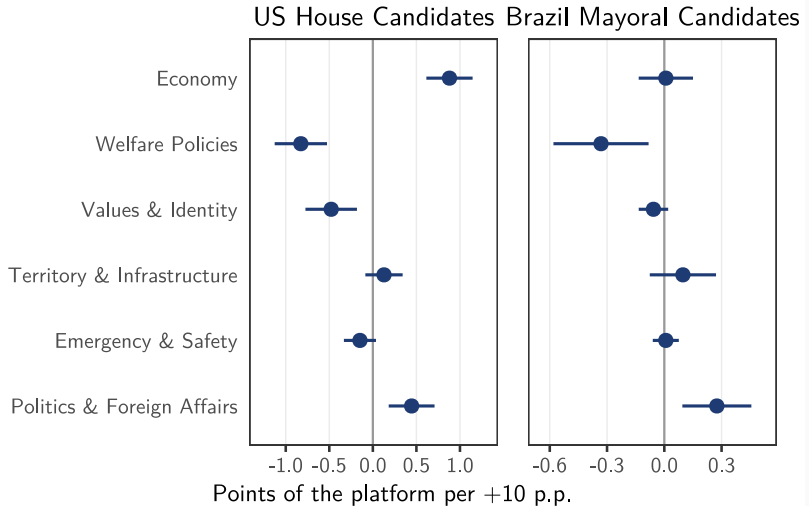
- Candidates strategically reposition to match the leaning of the district on presidential elections ([Meisels, 2026](#)).
- **Explanatory variable:** vote share of rightist presidential candidate.

Candidate Strategic Positioning

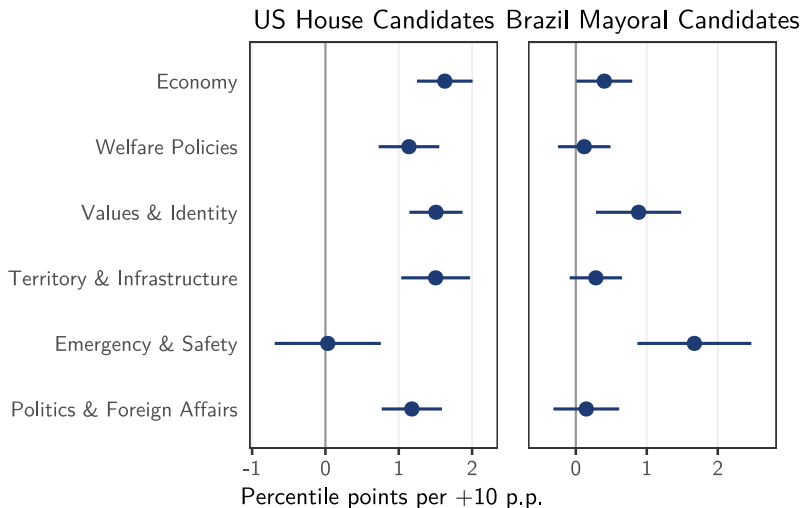
- Candidates strategically reposition to match the leaning of the district on presidential elections ([Meisels, 2026](#)).
- **Explanatory variable:** vote share of rightist presidential candidate.



Issue Salience Positioning



Issue Stance Positioning



- **Interpretability gain:** Fasti identifies which issues divide, by salience and by stance, so construct validity is easy to assess.

- **Interpretability gain:** Fasti identifies which issues divide, by salience and by stance, so construct validity is easy to assess.
- **LLM caveat:** minor prompt variations (party and name cues) move extreme candidates; implausible divides in Brazil, i.e., international affairs and immigration.

Conclusion

- **Interpretability gain:** Fasti identifies which issues divide, by salience and by stance, so construct validity is easy to assess.
- **LLM caveat:** minor prompt variations (party and name cues) move extreme candidates; implausible divides in Brazil, i.e., international affairs and immigration.
- Even with LLM advances, the IRT model will remain useful: interpretable parameters; quick estimation.

Thank you!

Jefferson Leal

University of Rochester

jefferson.leal@rochester.edu

jeff-leal.github.io

- Angelov, D. (2020). Top2Vec: Distributed Representations of Topics. arXiv preprint arXiv:2008.09470.
- Bleck, J. and van de Walle, N. (2012). Valence Issues in African Elections: Navigating Uncertainty and the Weight of the Past. *Comparative Political Studies*, 46(11):1394–1421.
- Bolognesi, B., Ribeiro, E., and Codato, A. (2023). A New Ideological Classification of Brazilian Political Parties. *Dados*, 66:e20210164.
- Bonica, A. (2014). Mapping the Ideological Marketplace. *American Journal of Political Science*, 58(2):367–386.
- Bucchianeri, P. (2020). Party Competition and Coalitional Stability: Evidence from American Local Government. *American Political Science Review*, 114(4):1055–1070.

- Burnham, M., Kahn, K., Wang, R. Y., and Peng, R. X. (2026). Political DEBATE: Efficient Zero-Shot and Few-Shot Classifiers for Political Text. *Political Analysis*, 34(3):329–343.
- Case, C. R. (2025). Measuring Strategic Positioning in Congressional Elections. *Journal of Politics*. Forthcoming.
- Clinton, J., Jackman, S., and Rivers, D. (2004). The Statistical Analysis of Roll Call Data. *American Political Science Review*, 98(2):355–370.
- Denny, M. J. and Spirling, A. (2018). Text Preprocessing for Unsupervised Learning: Why It Matters, When It Misleads, and What to Do about It. *Political Analysis*, 26(2):168–189.
- Godambe, V. P. (1960). An Optimum Property of Regular Maximum Likelihood Estimation. *The Annals of Mathematical Statistics*, 31(4):1208–1211.

- Grimmer, J. and Stewart, B. M. (2013). Text as Data: The Promise and Pitfalls of Automatic Content Analysis Methods for Political Texts. *Political Analysis*, 21(3):267–297.
- Grootendorst, M. (2022). BERTopic: Neural topic modeling with a class-based TF-IDF procedure. arXiv preprint arXiv:2203.05794.
- Kitschelt, H. (2007). Party Systems. In Boix, C. and Stokes, S. C., editors, *The Oxford Handbook of Comparative Politics*, pages 522–554. Oxford University Press, Oxford.
- Lauderdale, B. E. and Herzog, A. (2016). Measuring Political Positions from Legislative Speech. *Political Analysis*, 24(3):374–394.
- Laurer, M., van Atteveldt, W., Casas, A., and Welbers, K. (2024). Less Annotating, More Classifying: Addressing the Data Scarcity Issue of Supervised Machine Learning with Deep Transfer Learning and BERT-NLI. *Political Analysis*, 32(1):84–100.

- Laver, M. (2001). Position and Salience in the Policies of Political Actors. In Laver, M., editor, *Estimating the Policy Positions of Political Actors*, pages 66–75. Routledge, London.
- Laver, M., Benoit, K., and Garry, J. (2003). Extracting policy positions from political texts using words as data. *American Political Science Review*, 97(2):311–331.
- Lehmann, P., Franzmann, S., Al-Gaddooa, D., Burst, T., Ivanusch, C., Lewandowski, J., Regel, S., Riethmüller, F., and Zehnter, L. (2025). Manifesto Corpus. Version 2025-1. Berlin: WZB Berlin Social Science Center / Göttingen: Institute for Democracy Research (IfDem).
- Lowe, W., Benoit, K., Mikhaylov, S., and Laver, M. (2011). Scaling Policy Preferences from Coded Political Texts. *Legislative Studies Quarterly*, 36(1):123–155.

- McInnes, L., Healy, J., and Astels, S. (2017). hdbscan: Hierarchical density based clustering. *Journal of Open Source Software*, 2(11):205.
- McInnes, L., Healy, J., and Melville, J. (2018). UMAP: Uniform Manifold Approximation and Projection for Dimension Reduction. arXiv preprint arXiv:1802.03426.
- Meisels, M. (2026). Candidate Positions, Responsiveness, and Returns to Extremism. *The Journal of Politics*, 88(2):681–697.
- Merz, N., Regel, S., and Lewandowski, J. (2016). The Manifesto Corpus: A new resource for research on political parties and quantitative text analysis. *Research & Politics*, 3(2).
- Mikhaylov, S., Laver, M., and Benoit, K. (2012). Coder Reliability and Misclassification in the Human Coding of Party Manifestos. *Political Analysis*, 20(1):78–91.

- OpenAI (2025). Introducing GPT-4.1 in the API.
<https://openai.com/index/gpt-4-1/>.
- Porter, R., Case, C. R., and Treul, S. A. (2025). CampaignView, a database of policy platforms and biographical narratives for congressional candidates. *Scientific Data*, 12(1):1237.
- Reimers, N. and Gurevych, I. (2019). Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 3982–3992, Hong Kong.
- Rheault, L. and Cochrane, C. (2020). Word embeddings for the analysis of ideological dimensions in political texts. *Political Analysis*, 28(1):112–133.

- Slapin, J. B. and Proksch, S.-O. (2008). A scaling model for estimating time-series party positions from texts. *American Journal of Political Science*, 52(3):705–722.
- Tribunal Superior Eleitoral (2020). Candidatos: Eleições municipais 2020 (candidate microdata). Repositório de Dados Eleitorais, <https://dadosabertos.tse.jus.br/>.
- Varin, C., Reid, N., and Firth, D. (2011). An Overview of Composite Likelihood Methods. *Statistica Sinica*, 21(1):5–42.
- Wang, W., Wei, F., Dong, L., Bao, H., Yang, N., and Zhou, M. (2020). MiniLM: Deep Self-Attention Distillation for Task-Agnostic Compression of Pre-Trained Transformers. In *Advances in Neural Information Processing Systems*, volume 33, pages 5776–5788.
- White, H. (1982). Maximum Likelihood Estimation of Misspecified Models. *Econometrica*, 50(1):1–25.

Appendix Contents

- Related Work
- Classical Text Scaling
- Wordfish Model
- Wordfish on Brazil: Choices
- Corpora
- Sentence Segmentation
- Topic Modeling
- Topic Model Details
- Stance Labels
- CMP Crosswalk
- Stance by Transfer Learning
- Training Labels & Filter
- Hypotheses & Training
- Held-Out Performance
- Item Response Model
- Identification & Estimation
- Centering & Group Means
- Ideal Point Distributions
- Model Comparison
- Model Comparison: Within Party
- Strategic Positioning: All Measures
- Conclusion

Related Work

- **Unsupervised scaling:** Wordscores, Wordfish ([Laver et al., 2003](#); [Slapin and Proksch, 2008](#)); embedding-based scaling ([Rheault and Cochrane, 2020](#); [Case, 2025](#)).
- **Topic-wise scaling:** Wordshoal, a factor model over topic-specific positions ([Lauderdale and Herzog, 2016](#)).
- **Behavioral and expert benchmarks:** CFscores, DW-NOMINATE, expert placements ([Bonica, 2014](#); [Bolognesi et al., 2023](#)). External, not text-derived.
- **Stance detection with NLI:** entailment classifiers reach strong accuracy with modest training data ([Laurer et al., 2024](#); [Burnham et al., 2026](#)).
- **Platforms vs. roll calls:** what candidates say is related to, not the same as, how they vote ([Meisels, 2026](#)).

- Bag-of-words models e.g., Wordfish.
- **Assumption:** word choice reveals issue salience: 'taxes' read as right, 'environment' as left ([Grimmer and Stewart, 2013](#); [Lowe et al., 2011](#)).
- **Problem:** when sides share vocabulary but take opposite stances, the signal is too subtle to recover ([Lauderdale and Herzog, 2016](#); [Rheault and Cochrane, 2020](#)).
- On less-polarized text, many measurement tools are hard to interpret and have validity issues.

Wordfish Model

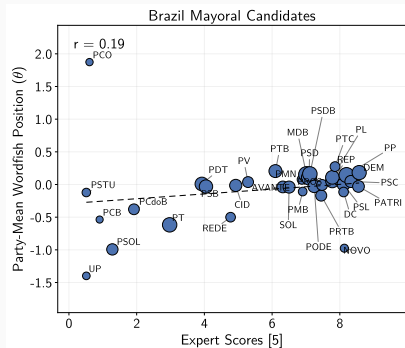
Feature w in document i , count Y_{iw} ; the document's location on the latent scale is θ_i (Slapin and Proksch, 2008):

$$Y_{iw} \sim \text{Poisson}(\lambda_{iw}), \quad \log \lambda_{iw} = \alpha_i + \psi_w + \beta_w \theta_i$$

- θ_i : the recovered scale (identified only up to location and scale).
- α_i : document length; ψ_w : feature baseline frequency.
- β_w : feature weight — pulls a document toward one pole.

Unsupervised: θ_i and every β_w are fit from counts alone, with no labels to fix which pole is which. Compare the joint model, where β_t and ρ_t load on supervised sentence codes.

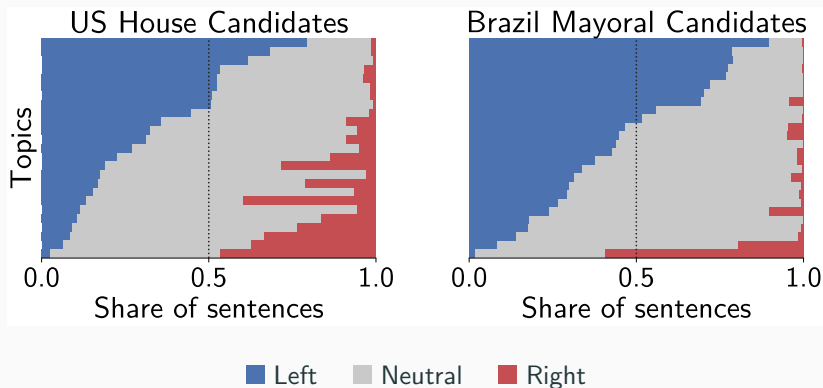
Wordfish on Brazil: Discretionary Choices



- Party means barely track the experts ($r = 0.19$).
- Discretionary choices have huge influence (Denny and Spirling, 2018): including bi-grams $r : 0.19 \rightarrow 0.73$.
- Algorithm only converges if one prunes rare terms (to $\sim 8,000$) and winsorizes counts of frequent terms.

Corpora

Corpora: US Congress campaign platforms, 2018–2022 primaries (CampaignView, [Porter et al., 2025](#)): 4,507 documents. Brazilian mayoral platforms, 2020 elections: 16,830 documents.



Sentence Segmentation

- One pipeline for both corpora: encoding repair, Unicode normalization, pre-trained spaCy sentence segmenter per language, list splitting, 150-word cap.
- Platform text is list-heavy; off-the-shelf segmentation alone over-splits.
- 439.3K chunks (US) and 3.10M (BR); median 18 and 15 words; median document breaks into 58 and 148 chunks.
- Every chunk is scored and counted. Nothing is discarded at this stage.
- Documents with <10 chunks are dropped from the scale: their mean stance is mechanically -1 , 0 , or $+1$. The scale is fit on 4,333 US and 16,674 Brazilian platforms.

Topic Modeling

- **Pipeline:** sentences $\rightarrow$ embeddings $\rightarrow$ topic model $\rightarrow$ stance $\rightarrow$ IRT.
- **Sentences:** spaCy + regex rules $\rightarrow$ chunks ≤ 150 words (median 18 and 15 words); 439.3K (US) and 3.10M (BR).
- **Contextual embeddings:** all-MiniLM-L6-v2 (US); paraphrase-multilingual-MiniLM-L12-v2 (BR) ([Wang et al., 2020](#); [Reimers and Gurevych, 2019](#)).
- **Topics:** dimensionality reduction via UMAP $\rightarrow$ HDBSCAN for clustering ([Grootendorst, 2022](#); [Angelov, 2020](#)).
- **US:** 96 subtopics. **Brazil:** 169 subtopics. Labeled by keywords and examples.
- **Supertopics:** subtopics grouped into 25 (US) and 26 (BR) classes of one codebook; the model estimates over those.

Topic Model

- **Embeddings:** 384 dimensions. all-MiniLM-L6-v2 (US); paraphrase-multilingual-MiniLM-L12-v2 (BR).
- **Reduction:** UMAP to 10 dimensions ([McInnes et al., 2018](#)).
- **Clustering:** HDBSCAN ([McInnes et al., 2017](#)): density-based, no preset topic count, explicit noise cluster.
- **US:** fit on the full corpus $\rightarrow$ 96 subtopics.
- **Brazil:** fit on 150 platforms per party (4,068 platforms; 770,770 chunks), remaining 2.3M chunks projected through the fitted reduction and cluster membership model $\rightarrow$ 169 subtopics.
- Noise chunks assigned to their highest-probability topic, so every sentence enters Y_{it} .
- Key terms by normalized class-based TF-IDF.

Stance Labels

Each sentence takes a stance: **Left**, **Right**, or **Neutral**, defined over two dimensions.

	Left	Right
Economy	State steers the economy: regulation, public services, redistribution, welfare, labour, environment	Market allocates: free enterprise, business incentives, free trade, fiscal orthodoxy, limited welfare
Society	Rights, equality, openness: civil rights, democracy & participation, minorities, immigration, peace	Nation, tradition, authority: patriotism, traditional morality, law & order, strong government

Neutral: no direction. Vague goals, valence and competence claims, anti-corruption, decentralization.

CMP Crosswalk: Economy

Left	Market Regulation (403), Economic Planning (404), Corporatism/Mixed Economy (405), Protectionism+ (406), Keynesian Demand Management (409), Controlled Economy (412), Nationalisation (413), Marxist Analysis (415), Anti-Growth/Sustainability+ (416.1–.2), Environmental Protection (501), Welfare State Expansion (504), Education Expansion (506), Labour Groups+ (701)
Right	Free Market Economy (401), Incentives+ (402), Protectionism– (407), Economic Growth+ (410), Economic Orthodoxy (414), Welfare State Limitation (505), Education Limitation (507), Labour Groups– (702)
Neutral	Economic Goals (408), Technology & Infrastructure+ (411), Culture+ (502), Agriculture & Farmers+ / – (703.1–.2)

+ / – abbreviate Positive/Negative; sub-codes collapse into their parent.

Codes conflating opposite poles (bare 201) are excluded, no guesses.

[Society](#)[Contents](#)[Back](#)

CMP Crosswalk: Society

Left Anti-Imperialism (103), Military— (105), Peace (106), Internationalism+ (107), Human Rights (201.2), Democracy (202), Pre-Democratic Elites— (305.5–.6), Equality+ (503), National Way of Life— (602), Traditional Morality— (604), Law & Order— (605.2), Multiculturalism+ (607), Underprivileged Minorities (705), Demographic Groups (706)

Right Military+ (104), Internationalism— (109), Freedom (201.1), Democracy: General— (202.2), Strong Government (305.3), Pre-Democratic Elites+ (305.4), National Way of Life+ (601), Traditional Morality+ (603), Law & Order+ (605), Civic Mindedness+ (606), Multiculturalism— (608)

Neutral Foreign Special Relationships+/- (101–102), EU+/- (108, 110), Constitutionalism+/- (203–204), Decentralization (301), Centralisation (302), Political Corruption (304), Political Competence (305.1–.2), Bottom-Up Activism (606.2), Middle Class & Professionals (704)

Stance by Transfer Learning

- **Labels:** Transfer learning from Manifesto Project (sentence-level labels); 12 countries, 9 languages ([Lehmann et al., 2025](#); [Merz et al., 2016](#)).
- **Filter:** Keep only sentences where GPT-4.1 agrees with the human coder ([Burnham et al., 2026](#)) to mitigate low intercoder reliability ([Mikhaylov et al., 2012](#)).
- **Classifier:** fine-tuned multilingual natural language inference (NLI) model (requires small training data) ([Laurer et al., 2024](#); [Burnham et al., 2026](#)).
 - Four tasks: dimension (social or economic), economic stance, social stance, general stance.
 - **Held-out balanced accuracy 0.83** over three classes.

Stance: Training Labels & Filter

- Stratified sample: up to 1,000 sentences per language–dimension–stance cell; full sentences of ≥ 8 words; **48,982 sentences**.
- Every sentence recoded by GPT-4.1 ([OpenAI, 2025](#)). Keep only concordant sentences: a filter, not an adjudication.

Task		n	Agreement	κ
T1	Dimension (Economy/Society)	48,982	78.8%	0.58
T2	Economic pole	22,006	51.5%	0.27
T3	Social pole	16,586	59.4%	0.40
T4	General stance (production)	48,982	52.6%	0.28

The disagreement is directional: GPT-4.1 recovers 68.2% of the human coder's Left sentences but only 43.5% of the Right ones. Concordant training set: **25,761** stance-labeled sentences.

Stance: Hypotheses & Training

- **NLI framing:** premise (the sentence) + hypothesis $\rightarrow$ entailment. Production hypotheses: “This text expresses a {leftist, rightist, neutral} political stance.”
- Four tasks, eleven hypotheses: dimension, economic pole, social pole, general stance.
- **Base model:** bge-m3-zeroshot-v2.0 ([Laurer et al., 2024](#)), binary entailment head.
- Each sentence yields an entailing and a non-entailing pair, in its language and in English translation: **394,087 pairs** from 43,173 sentences.
- Split grouped by *manifesto* (70/15/15): evaluation is on parties and documents never seen in training.
- Early stopping on validation Matthews correlation; one A100 trains it in under an hour.

Stance: Held-Out Performance

Task	n	MCC	Balanced accuracy
T1: Dimension	6,176	0.88	0.94
T2: Economic pole	1,678	0.76	0.84
T3: Social pole	1,626	0.77	0.84
T4: General stance	4,172	0.75	0.83

T4 balanced accuracy by language

Catalan	0.86	French	0.81
Swedish	0.85	German	0.79
Spanish	0.85	Galician	0.76
Italian	0.84	English	0.82
		Portuguese	0.82

Test manifestos never seen in training. Per-class F1: Left 0.86, Right 0.79, Neutral 0.85.

Item Response Model

Document i , topic t : Y_{it} = chunk count ($M_i = \sum_t Y_{it}$), $S_{it} \in [-1, 1]$ = mean stance, θ_i = ideal point.

Stance:

$$S_{it} \sim \mathcal{N}\left(\alpha_t + \beta_t \theta_i, \frac{\sigma^2}{Y_{it}}\right)$$

Counts:

$$Y_i \sim \text{Multinomial}(M_i, \pi_i), \quad \pi_{it} = \exp\{\eta_{it}\} / \sum_s \exp\{\eta_{is}\}$$
$$\eta_{it} = \psi_t + \rho_t \theta_i + \kappa_t(\theta_i^2 - 1)$$

- α_t, ψ_t : topic's baseline lean and baseline salience. Consensus is absorbed here.
- β_t : **stance gap**. Which issues split the sides by stance.
- ρ_t : **salience gap**. Which issues each side relatively emphasizes.
- κ_t : curvature. Allows for non-monotone salience.

Location, scale, and reflection are conventions, not findings:

$$\sum_i \theta_i = 0, \quad \frac{1}{T} \sum_t \beta_t^2 = 1, \quad \bar{\theta}_{\text{right}} > 0$$

- **Stage 1:** calibrate parameters from the stance equation $(\alpha_t, \beta_t, \theta_i, \sigma^2)$.
- **Stage 2:** jointly fit both equations (Godambe-corrected composite likelihood) ([Godambe, 1960](#); [Varin et al., 2011](#)).
- Composite likelihood $\rightarrow$ robust standard errors, documents as clusters.

Estimation: Objective

Stage 1. Pilot fit of the stance equation, by alternating least squares:

$$(\hat{\alpha}, \hat{\beta}, \theta^S) = \arg \min \sum_{i, t: Y_{it} > 0} Y_{it} (S_{it} - \alpha_t - \beta_t \theta_i)^2$$

Calibrate and freeze: $\hat{\sigma}^2$, an overdispersion factor d_i , and a Godambe factor c_Y .

Stage 2. One calibrated objective over θ and all item parameters:

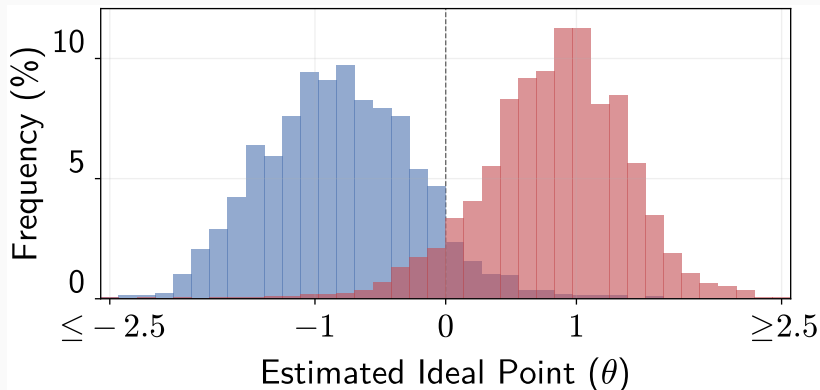
$$\mathcal{L} = \sum_i \left[\frac{\ell_i^S}{2\hat{\sigma}^2} + \frac{\ell_i^Y}{d_i c_Y} + \frac{\theta_i^2}{2} \right] + \text{pen}(\kappa)$$

- d_i discounts documents that restate the same point: platform text repeats itself.
- c_Y matches the count block's weight to its information about θ_i , not its nominal likelihood ([Godambe, 1960](#); [Varin et al., 2011](#)).
- Block coordinate descent; robust SEs from the Godambe information ([White, 1982](#)).

Centering and Group Means

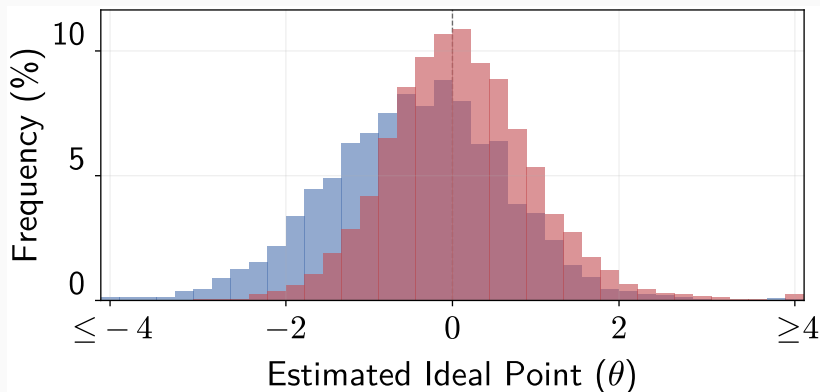
- $\sum_i \theta_i = 0$ is a normalization. With two groups exhausting the sample, $n_1 \bar{\theta}_1 = -n_2 \bar{\theta}_2$.
- **US:** 2,187 Democratic and 2,146 Republican platforms; fitted means -0.209 and $+0.213$. Their ratio, 1.019 , is the ratio of the counts.
- No feature of the text produced that symmetry: either mean determines the other.
- **What carries information:** the separation between groups, the within-group spread, the shape of the distribution. Not the individual means ([Lowe et al., 2011](#); [Clinton et al., 2004](#)).

Ideal Points: United States



■ Democrats ■ Republicans

Ideal Points: Brazil



■ Left Parties ■ Other Parties

Model Comparison

Model	United States		Brazil
	CFScore	DW-NOMINATE	Expert score
<i>Unsupervised scaling, bag-of-words</i>			
Correspondence analysis, unigrams	0.53	0.70	0.64
Correspondence analysis, bigrams	0.65	0.74	0.54
Wordfish, unigrams	0.72	0.80	0.19
Wordfish, bigrams	0.76	0.81	0.73
<i>Document embeddings</i>			
Doc2Vec + PCA, unigrams	0.19	0.47	0.70
Doc2Vec + PCA, bigrams	0.21	0.49	0.69
Pre-trained Embeddings + PCA	0.56	0.57	0.69
<i>Large language models</i>			
Ask and average, GPT-4o mini	0.87	0.89	0.82
Ask and average, GPT-5.4 mini	0.85	0.88	0.82
<i>This paper</i>			
Fasti	0.81	0.85	0.86
<i>n</i>	3,788	1,111	33

Pearson correlation of θ with each benchmark. Bold: best per column.

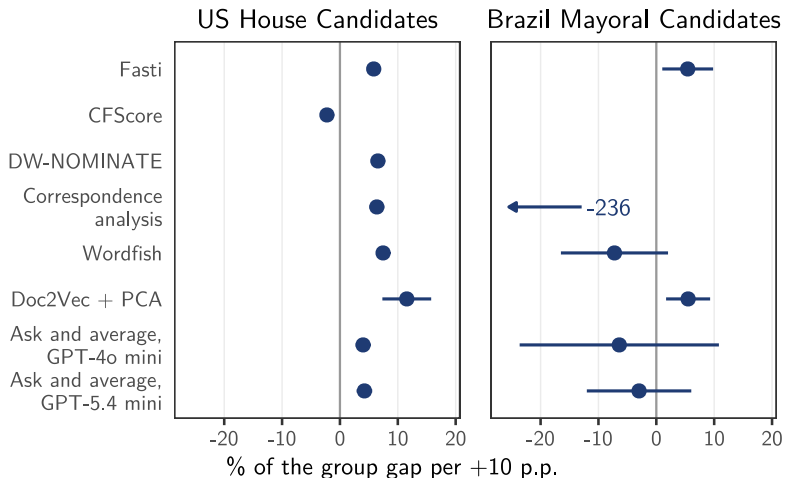
Model Comparison: Within Party

Model	CFScore		DW-NOMINATE	
	Dem.	Rep.	Dem.	Rep.
<i>Unsupervised scaling, bag-of-words</i>				
Correspondence analysis, unigrams	-0.17	0.13	0.10	0.40
Correspondence analysis, bigrams	-0.17	0.26	0.10	0.41
Wordfish, unigrams	-0.09	0.25	0.16	0.38
Wordfish, bigrams	-0.03	0.23	0.22	0.34
<i>Document embeddings</i>				
Doc2Vec + PCA, unigrams	-0.23	0.12	0.02	0.26
Doc2Vec + PCA, bigrams	-0.23	0.12	0.00	0.26
Pre-trained Embeddings + PCA	-0.08	0.23	-0.02	0.38
<i>Large language models</i>				
Ask and average, GPT-4o mini	0.17	0.20	0.20	0.26
Ask and average, GPT-5.4 mini	0.22	0.21	0.19	0.27
<i>This paper</i>				
Fasti	0.10	0.20	0.29	0.28
<i>n</i>	2,020	1,768	576	535

Bag-of-words and embedding baselines order Democrats against CFScore.

[Contents](#)[All candidates](#)[Back](#)

Strategic Positioning: All Measures



A row outside the axis is an arrow through the edge, with its estimate.

Conclusion

- **Main result:** the method recovers interpretable ideal points from text even where party language is *not* polarized (i.e., Brazil).
- **US:** parties use distinct stances and framings, reflecting cultural wars and divergent economic positions.
- **Brazil:** corpus is dominated by support for public services (or neutral). Left: political institutions, minorities, social assistance, values; Right: urban management, technology, infrastructure, public safety, business.
- Within Brazilian parties, candidates vary widely in ideology, yet party means track the national ordering experts perceive.
- Topic + stance outperform embeddings + PCA ([Rheault and Cochrane, 2020](#)); extracting topic and stance from the embeddings greatly improves the estimates.
- No target-domain (US/BR) labels needed (given transfer learning of stance); scalable on modest GPU.