

Multilingual Power and Ideology Identification in the Parliament: a Reference Dataset and Simple Baselines

Çağrı Çöltekin¹, Matyáš Kopp², Katja Meden^{3,5},
Vaidas Morkevicius⁴, Nikola Ljubešić³, Tomaž Erjavec³

¹University of Tübingen, Tübingen, Germany, ²Charles University, Prague, Czech Republic,

³Jožef Stefan Institute, Ljubljana, Slovenia, ⁴Kaunas University of Technology, Kaunas, Lithuania,

⁵Jožef Stefan International Postgraduate School, Slovenia,

ccoltekin@sfs.uni-tuebingen.de, kopp@ufal.mff.cuni.cz, katja.meden@ijs.si,

vaidas.morkevicius@ktu.lt, nikola.ljubestic@ijs.si, tomaz.erjavec@ijs.si

Abstract

We introduce a dataset on political orientation and power position identification. The dataset is derived from ParlaMint, a set of comparable corpora of transcribed parliamentary speeches from 29 national and regional parliaments. We introduce the dataset, provide the reasoning behind some of the choices during its creation, present statistics on the dataset, and, using a simple classifier, some baseline results on predicting political orientation on the left-to-right axis, and on power position identification, i.e., distinguishing between the speeches delivered by governing coalition party members from those of opposition party members.

Keywords: ideology, power, parliamentary corpus, ParlaMint

1. Introduction

Parliaments are one of the most important institutions in modern democratic states where issues with high societal impact are discussed. The decisions made in a national parliament affect the citizens of its country on fundamental aspects of their life. The societal importance of parliamentary discourse requires a better understanding and analysis of parliamentary debates. As a result, there has been a recent increase in the number of resources (Fišer and Lenardič, 2018; Lenardič and Fišer, 2023) and (computational) linguistic analyses of parliamentary debates (see Glavaš et al., 2019; Abercrombie and Batista-Navarro, 2020, for recent reviews). The impact of the decisions made in a parliament often goes beyond their borders, and may even have global effects. Hence, comparative studies of parliamentary debates across countries and in multiple languages is also important.

The dataset described here is derived from the ParlaMint corpora, a collection of comparable corpora of transcribed parliamentary speeches from 29 national and regional parliaments, covering at least the period from 2015 to 2022 (Erjavec et al., 2022). The dataset is prepared for a shared task on two important aspects of a political discourse, *political orientation* and *power* (Kiesel et al., 2024).¹ Although a simplification, political orientation on the left-to-right spectrum

has been one of the defining properties of political ideology (Arian and Shamir, 1983; Vegetti and Širinić, 2019). Power is another factor that shapes the political discourse (van Dijk, 2008; Fairclough, 2013a,b). Despite its central role in critical discourse analysis, to the best of our knowledge, power was not studied computationally earlier.² We provide a reference dataset of parliamentary speeches for both tasks, which we expect to be instrumental for quantitative and computational studies on ideology and power in parliamentary debates beyond the present shared task as well.

Both tasks are formulated as binary classification tasks. For the power position identification task, this choice is mostly straightforward, as the distinction we want to make is between the speeches delivered by governing party members and those given by opposition party members.

Classifying political orientation is more complex, as it can be expressed in many ways. In fact, ParlaMint provides annotations from two sources (Erjavec et al., 2023b): Wikipedia and the Chapel Hill Expert Survey Europe (CHES, Jolly et al., 2022). Wikipedia classifies the political orientations of parties into 13 categories on the left-to-right spectrum, as well as five other values that do not fit into this axis (e.g., ‘Big Tent’, or ‘Single Issue Politics’ values). Conversely, CHES gives political orienta-

[in-parliamentary-debates.html](#).

²Our definition of power for the present data set is also simplified. As suggested by an anonymous reviewer, other power roles, such as being a (shadow) cabinet member, or the role in the party may manifest differently in the speech. We leave such aspect of power in speech for future research.

¹Further practical information about the shared task can be found on the shared task web page at <https://touche.webis.de/clef24/touche24-web/ideology-and-power-identification>

tion along a large number of dimensions (85 in total, e.g., stance towards European integration, but also the general left-to-right position of a party), with the numeric values based on averaged scores of expert surveys. For the left-to-right position experts assigned a numeric score between 0 to 10 (far left to far right) based on a party’s general ideological stance. Not all parties have political orientation annotations in ParlaMint, but the coverage of the Wikipedia annotations is more comprehensive than that of the CHES annotations. As a result, we use orientation values from Wikipedia.

To facilitate graded predictions on the left-to-right scale, we use labels 0 for left, and 1 for right-wing parties. We mark Wikipedia categories from ‘far-left’ (FL) to ‘centre to centre-left’ (CCL) as *left*, and those from ‘far-right’ (FR) to ‘centre to centre-right’ (CCR) as *right*. We exclude the speeches from the members of the parties marked as centre and parties whose orientation does not fit into the left-to-right continuum.

For both tasks, the main challenge in the creation of a dataset is to minimize the effects of covariates. Even though the instances to classify are speeches, the annotations are based on the party membership of the speaker. As a result, underlying variables like party membership, or speaker identity perfectly covary with ideology and power in most cases. The sampling procedure described in Section 2 below aims to reduce these correlations, and encourage systems trained on the data to generalize to the particular task, rather than predictions based on easier-to-guess covariates.

ParlaMint is a multilingual dataset of transcribed speeches delivered in different regional and national parliaments. As a result, it also offers opportunities to investigate similarities and differences of ideology and power in varying cultures and parliamentary traditions, as well as their reflection in different languages. Even though the shared task does not offer a cross-lingual evaluation track, the uniformly encoded data allows participants to exploit ‘universal’ aspects of ideology and power through, for example, transfer learning. To encourage participation in multiple languages, and help participants build (simple) multilingual classifiers easily, we also include automatic English translations of the speeches.

Our aim in this paper is to describe the process and rationale behind the dataset construction, as well as providing an overview of the resulting data. We also describe a trivial baseline and the results of experiments with this baseline.

2. Data

The data is a subset of ParlaMint version 4.0 (Erjavec et al., 2023a). For the shared task, we

split the data into training and test sets (without a fixed validation set), and share them via <https://zenodo.org/records/10450640>. We also provide English translations provided in the ParlaMint distribution (Kuzman et al., 2023). The main motivation for the subsampling is to reduce the effects of covariates explained above. Furthermore, since ParlaMint contains over 1.2 billion words, and more than 7.7 million speeches (more correctly ‘utterances’ in ParlaMint TEI annotations), sampling also results in a more manageable dataset for machine-learning experiments, promoting inclusion of participants without access to high-performance computing facilities.

Before sampling the speeches, we join the utterances by the same speaker when they were interrupted by a single utterance of another speaker, and we filter out speeches that are shorter than 500 characters, and longer than 20 000 characters. The former is intended for the inclusion of the interrupted speeches as a whole.³ The latter, filtering by size, removes short interruptions and very long speeches. On average, the lengths of the selected speeches are between 200 and 1 000 words, approximately corresponding to speeches of 2 to 10 minutes. The utterances of the session chairs, which are typically about procedural matters, are always filtered out.

The only preprocessing steps we apply are replacing the party names or abbreviations as listed in ParlaMint with a placeholder <PARTY>, and using a <p> tag to indicate paragraph boundaries in the original transcripts. Masking the party references eliminates some trivial cues, as in ‘I am speaking on behalf of <PARTY>’. We only replace the party names and abbreviations as given in ParlaMint metadata, which do not cover some of the alternative names or abbreviations of the parties, as well as (consistent) mistranslations in the automatically translated texts. We leave the rest of the named entities intact. Even though (stance towards) some of the named entities may also provide strong cues for power and ideology, many of these cues will be legitimate, and we expect the models to discover and make use of them (e.g., the stance towards a particular event, like Brexit, may genuinely stem from a speakers’ relation with the government or their political orientation). Future releases of the data may improve on eliminating the obvious cues for power or ideology.

We also include the sex of the speaker, an anonymised speaker ID, and automatic translation to English in the training data. The gender information in ParlaMint was collected from var-

³It is common for the speeches to be interrupted by the chair, often asking the speaker to finish in the allotted time. Unauthorized interruptions from the audience are also common.

	Orientation						Power					
	Training			Test			Training			Test		
	n	L%	tokens	n	L%	tokens	n	O%	tokens	n	O%	tokens
Austria (AT)	7 879	32.6	535.4	2 002	44.7	566.6	15 971	58.8	568.1	2 181	49.0	598.5
Bosnia and Herzegovina (BA)	1 301	20.9	375.4	2 014	28.9	348.2	2 531	16.8	351.5	1 992	16.9	355.0
Belgium (BE)	2 276	32.1	403.9	2 018	38.2	378.4	4 765	47.4	397.1	1 973	47.4	398.2
Bulgaria (BG)	3 907	32.3	447.9	2 006	36.0	444.8	6 699	52.8	444.6	1 981	46.1	456.9
Czechia (CZ)	4 137	39.0	356.9	2 002	18.8	386.9	6 744	47.8	376.2	1 965	42.9	406.5
Denmark (DK)	3 069	57.1	457.2	2 015	56.6	465.7	5 493	37.2	498.8	1 971	47.4	529.7
Estonia (EE)	2 595	36.4	243.6	2 012	38.9	247.5	-	-	-	-	-	-
Spain (ES)	4 770	44.9	938.2	2 003	53.8	956.3	7 198	29.3	935.7	1 930	40.9	960.5
Catalonia (ES-CT)	2 077	46.6	915.2	2 007	47.5	921.0	1 525	34.8	896.0	1 999	35.3	904.1
Galicia (ES-GA)	943	54.1	1 072.1	2 010	58.2	1 144.2	953	42.5	1 138.0	2 000	43.5	1 164.0
Basque Country (ES-PV)	-	-	-	-	-	-	1 031	43.7	962.6	1 989	46.3	981.9
Finland (FI)	1 179	42.7	233.2	2 001	45.5	219.8	6 111	55.4	227.3	1 986	49.6	219.3
France (FR)	3 618	30.2	275.3	2 002	28.2	292.8	9 813	63.0	272.3	1 996	66.5	275.3
Great Britain (GB)	24 239	48.8	438.5	2 017	44.7	465.9	33 257	43.6	455.0	1 996	31.9	485.7
Greece (GR)	5 639	46.9	959.8	2 013	56.7	959.7	6 389	37.3	971.0	1 972	42.8	966.4
Croatia (HR)	8 322	22.8	489.7	2 016	26.9	504.2	10 741	60.3	503.9	1 989	58.8	525.8
Hungary (HU)	2 935	24.2	581.3	2 020	24.0	633.0	2 597	59.1	598.8	2 000	57.7	585.7
Iceland (IS)	536	48.0	470.0	2 015	38.3	552.5	-	-	-	-	-	-
Italy (IT)	3 367	38.3	696.5	2 014	45.8	707.4	7 848	62.5	671.7	1 971	56.8	704.5
Latvia (LV)	798	21.3	357.9	2 008	19.5	303.9	1 410	67.0	317.5	1 990	70.5	303.3
The Netherlands (NL)	5 657	38.4	502.5	2 001	37.8	473.0	7 906	58.5	484.5	1 986	59.4	500.7
Norway (NO)	10 998	50.4	457.1	2 009	40.8	475.7	-	-	-	-	-	-
Poland (PL)	5 489	11.1	356.4	2 014	16.9	359.6	9 705	45.2	329.8	2 000	46.3	340.1
Portugal (PT)	3 464	57.7	459.3	2 001	56.1	464.9	7 692	58.7	458.6	1 958	43.2	451.9
Serbia (RS)	9 914	16.1	652.9	2 015	14.1	594.5	15 114	72.9	650.4	1 990	65.7	659.2
Sweden (SE)	8 425	46.3	675.2	2 011	47.4	702.1	-	-	-	-	-	-
Slovenia (SI)	2 726	73.4	516.4	2 002	63.5	519.5	9 040	62.5	533.6	2 014	49.7	526.7
Turkey (TR)	16 138	41.8	410.3	2 008	45.7	413.7	17 384	48.6	418.5	1 990	44.5	430.3
Ukraine (UA)	2 545	16.2	232.3	2 001	14.8	242.4	11 324	68.8	224.5	2 182	35.6	233.3

Table 1: Statistics of the dataset. For each dataset, the number of speeches (n), the class imbalance (L% – the percentage of *left* for orientation, O% – the percentage of *opposition* for power), and the average number of tokens are reported.

ious sources, typically from the information provided on the web pages of the parliaments, or from Wikipedia, while in a small number of cases, the gender is unknown. Similarly, the machine translations are also not available in a small number of instances, mostly due to technical problems. The motivation for including speaker ID is to provide informed ways of dividing the available data as training and validation sets. The speaker ID is not included in the test set.

Sampling For ideal datasets for both tasks, we would need a large variation with respect to political party affiliations and speaker identities. For example, we would want multiple disjoint left-wing and right-wing political parties to be present in the training set and the test set so that the models could be evaluated for their ability to predict political orientation without relying on party affiliation. However, the nature of the ParlaMint data (in fact, any realistic corpus of parliamentary debates) prevents having such a dataset. For many parliaments, the number of political parties of a particular orientation is limited to a small number. For the power identification tasks, this is even more severe since a single party or only a few parties are

in power in some countries throughout the time period covered in ParlaMint.

As a trade-off between data size, and for reducing the effect of covariates, we opt for a speaker-based sampling. First, to discourage, to some extent, the classifiers from relying on author identification, we sample maximally 20 speeches of a single speaker. This is also important for introducing variation into the dataset, as the number of speeches from each speaker follows a power-law distribution. While a small number of speakers tend to deliver most of the speeches, e.g., party or party group leaders, most speakers have relatively few speeches. The distribution of speeches or speakers to include in training and test sets is also important for proper evaluation. For the ideology task, the set of speakers in the training and test sets are disjoint. For a reasonably accurate evaluation, we set the test set size to 2 000 instances (about 100 to 200 speakers depending on the individual corpus and the task). Despite multiple speeches from each speaker, due to missing annotations and the lack of diversity of orientation in some parliaments, the disjoint training/test constraint above results in a small number of training instances, leaving a small number of instances in

the training set for some of the parliaments.

Ideally, power identification requires a different constraint. That is, the same speaker should be present in both training and test sets such that speeches from one set should be when the speaker was in power, and the other set should contain the speeches while the same speaker is part of the opposition. This constraint is too difficult, or impossible, to satisfy for many parliaments in the ParlaMint data. For example, in Poland, only a single party is in power throughout the period covered by the corpus. Similarly, even when there is some variation, only a small number of speakers often serve both in governing coalitions and opposition. As a result, we use a best-effort train–test split, where if possible, we make sure that the speakers in the test set are also available in the training set with the opposite power role.⁴ Otherwise, we randomly sample more speakers to obtain approximately 2 000 instances in the test set. Political systems in some countries do not have a formal coalition–opposition distinction. As a result, we leave these parliaments out of the dataset.

Statistics The procedure described above results in training sets from 28 parliaments for the ideology identification task, and 25 parliaments for the power identification task. Table 1 provides some statistics on the training and test datasets. In general, there is a varying class imbalance in both datasets, but class distribution and speech lengths between training and test sets are similar. For some parliaments, the sampling procedure results in rather small training sets. Better classification of these datasets may be achieved by techniques like cross-lingual transfer and data augmentation.

3. Baselines

The main purpose of this paper is to introduce the dataset. However, we also report results from a simple baseline which is provided for the shared task. The baseline uses TF-IDF weighted character n-gram features with a simple logistic regression classifier. The motivation for such a simple baseline is twofold. First, since it will be used as the baseline for the shared task, a competitive baseline may intimidate some of the potential participants, particularly students and early researchers. Second, since the baseline only uses ‘surface’ features, with no claim of ‘language understanding’, it also provides initial data about how much of ‘the politics is about the words’.

Table 2 presents the F1-scores of the baseline for both tasks and for all parliaments. Most scores

	Orientation		Power	
	dev	test	dev	test
AT	59.1	51.9	68.5	65.0
BA	42.4	41.6	46.0	45.9
BE	55.6	56.7	58.3	63.4
BG	53.7	53.7	61.8	64.7
CZ	54.0	51.1	59.0	62.0
DK	50.9	54.0	51.7	53.4
EE	47.5	47.4	-	-
ES	72.1	71.7	61.2	65.0
ES-CT	72.8	66.4	68.6	76.7
ES-GA	62.4	70.5	74.3	70.7
ES-PV	-	-	66.3	68.9
FI	59.4	52.6	55.9	52.1
FR	43.9	45.0	64.1	66.1
GB	75.9	74.9	74.4	70.9
GR	72.5	75.2	66.9	64.0
HR	43.8	43.2	60.2	59.4
HU	56.2	55.8	81.8	84.9
IS	41.6	46.2	-	-
IT	57.3	50.9	47.0	43.9
LV	42.8	44.6	42.0	52.3
NL	51.4	54.4	60.9	64.5
NO	60.9	63.0	-	-
PL	46.4	45.4	74.6	75.6
PT	61.7	63.7	67.5	63.4
RS	47.9	51.6	69.7	62.7
SE	75.5	75.5	-	-
SI	44.5	40.7	53.1	53.7
TR	85.8	83.6	84.4	81.9
UA	56.7	58.9	59.4	45.4

Table 2: Macro-averaged F1-scores of the baseline on (dev)elopment and test sets on all development and test sets. All scores are averages of five random splits of the provided training data as 80 % for training and 20 % for validation. The scores above were obtained without any hyperparameter tuning.

are better than a random baseline (which would result in a 50 % F1-score). Most of the lower scores are the result of relatively high precision and low recall,⁵ clearly showing the lack of hyperparameter tuning. The mild correlation between the F1-scores and the training set size (0.53 and 0.36 on orientation and power detection tasks respectively) and weak but significant correlation of the class imbalance and the scores (−0.21 and −0.16 on orientation and power detection tasks respectively) also indicate that the data size and class imbalance are important factors for the success of the present classifier. However, these are not the only sources of difficulty. Despite relatively

⁴The data from only three parliaments (AT, SI, UA) satisfy this constraint, while there are no speakers that changed their roles in ES-GA, HU and PL.

⁵Since F1-score favours similar precision and recall values.

large datasets, for example, AT and NO are classified rather poorly for political orientation (and also the F1-score drops substantially in the test set compared to the development set), which may be because of better separation of speakers across training and test sets. On the other hand, the success of the baseline on both tasks on TR is unlikely to be explainable by the size and the class imbalance. One can perhaps relate these to political polarization, rather than the technical reasons we list above.⁶

4. Conclusions

The paper presents a dataset derived from the ParlaMint corpora, meant for studying automatic methods for detecting political orientation and power position in parliamentary debates. We believe it could be a valuable resource for studying these phenomena and other aspects of political discourse in multiple political and parliamentary cultures/traditions, and in multiple languages. Since measuring power and ideology on an individual basis is difficult, we use the well-known sources of party orientation and power position information to label individual speeches. This introduces some strong covariates of the ideology and power in any dataset that is derived from existing resources. Instead of a more restrictive setting where covariates are more strictly eliminated, we opted for a more inclusive dataset of including many parliaments and languages. We intend to improve the existing dataset by increasing its coverage and quality and by adding more metadata.

5. Limitations

The orientation and power based on party affiliation may not always reflect the individuals' positions at the time of their speeches. However, this is unlikely to be resolved easily without restricting the number of speakers drastically. A possible solution, as suggested by an anonymous reviewer, is to do manual annotations of the individual politicians by the experts, which would definitely be costly, and may also have its own limitations, such as changing positions in time.

We did not include the centre even though it clearly falls within the left–right spectrum of political orientation. This decision was motivated by simplicity. The inclusion of a centre in a binary classification scheme is not trivial, and not all parliamentary corpora include parties annotated as centre. For the future, multi-class classification, or

a form of ordinal regression/classification may be interesting alternatives against this limitation.

In the current version of the data, some procedural aspects of speech may also provide trivial, unwanted, cues for power and orientation. More rigorous identification and elimination of these cues in a big multilingual corpus is a difficult undertaking, that we leave for a potential new version of the corpus.

6. Acknowledgements

This work has been supported by CLARIN ERIC, ParlaMint: Towards Comparable Parliamentary Corpora.

7. Bibliographical References

- Gavin Abercrombie and Riza Batista-Navarro. 2020. Sentiment and position-taking analysis of parliamentary debates: a systematic literature review. *Journal of Computational Social Science*, 3(1):245–270.
- Asher Arian and Michal Shamir. 1983. The primarily political functions of the left-right continuum. *Comparative politics*, 15(2):139–158.
- Tomaž Erjavec, Matyáš Kopp, Maciej Ogrodniczuk, Petya Osenova, Manex Agirrezabal, Tommaso Agnoloni, José Aires, Monica Albini, Jon Alkorta, Iván Antiba-Cartazo, Ekain Arrieta, Mario Barcala, Daniel Bardanca, Starkaður Barkarson, Roberto Bartolini, Roberto Battistoni, Nuria Bel, Maria del Mar Bonet Ramos, María Calzada Pérez, Aida Cardoso, Çağrı Çöltekin, Matthew Coole, Roberts Dargis, Ruben de Libano, Griet Depoorter, Sascha Diwersy, Réka Dodé, Kike Fernandez, Elisa Fernández Rei, Francesca Frontini, Marcos Garcia, Noelia García Díaz, Pedro García Louzao, Maria Gavrilidou, Dimitris Gkoumas, Ilko Grigorov, Vladislava Grigorova, Dorte Haltrup Hansen, Mikel Iruskietia, Johan Jarlbrink, Kinga Jelencsik-Mátyus, Bart Jongejan, Neeme Kahusk, Martin Kirnbauer, Anna Kryvenko, Noémi Ligeti-Nagy, Nikola Ljubešić, Giancarlo Luxardo, Carmen Magariños, Måns Magnusson, Carlo Marchetti, Maarten Marx, Katja Meden, Amália Mendes, Michal Mochtak, Martin Mölder, Simonetta Montemagni, Costanza Navarretta, Bartłomiej Nitoń, Fredrik Mohammadi Norén, Amanda Nwadukwe, Mihael Ojsteršek, Andrej Pančur, Vassilis Papavasiliou, Rui Pereira, María Pérez Lago, Stelios Piperidis, Hannes Pirker, Marilina Pisani, Henk van der Pol, Prokopis Prokopidis, Valeria

⁶A proper investigation of this is beyond the scope of the current paper. Hence this statement should only be taken as a potential future direction for research.

- Quochi, Paul Rayson, Xosé Luís Regueira, Michał Rudolf, Manuela Ruisi, Peter Rupnik, Daniel Schopper, Kiril Simov, Laura Sinikallio, Jure Skubic, Lars Magne Tunland, Jouni Tuominen, Ruben van Heusden, Zsófia Varga, Marta Vázquez Abuín, Giulia Venturi, Adrián Vidal Miguéns, Kadri Vider, Ainhoa Vivel Couso, Adina Ioana Vladu, Tanja Wissik, Väinö Yrjänäinen, Rodolfo Zevallos, and Darja Fišer. 2023a. *Multilingual comparable corpora of parliamentary debates ParlaMint 4.0*. Slovenian language resource repository CLARIN.SI.
- Tomaž Erjavec, Maciej Ogrodniczuk, Petya Osenova, Nikola Ljubešić, Kiril Simov, Vladislava Grigorova, Michał Rudolf, Andrej Pančur, Matyáš Kopp, Starkaður Barkarson, Steinþór Steingrímsson, Henk van der Pol, Griet Depoorter, Jesse de Does, Bart Jongejan, Dorte Haltrup Hansen, Costanza Navarretta, María Calzada Pérez, Luciana D. de Macedo, Ruben van Heusden, Maarten Marx, Çağrı Çöltekin, Matthew Coole, Tommaso Agnoloni, Francesca Frontini, Simonetta Montemagni, Valeria Quochi, Giulia Venturi, Manuela Ruisi, Carlo Marchetti, Roberto Battistoni, Miklós Sebők, Orsolya Ring, Roberts Dargis, Andrius Utka, Mindaugas Petkevičius, Monika Briedienė, Tomas Krilavičius, Vaidas Morkevičius, Sascha Diwersy, Giancarlo Luxardo, and Paul Rayson. 2021. *Multilingual comparable corpora of parliamentary debates ParlaMint 2.1*. Slovenian language resource repository CLARIN.SI.
- Tomaž Erjavec, Katja Meden, and Jure Skubic. 2023b. Adding political orientation metadata to ParlaMint corpora. In *CLARIN annual conference 2023, book of abstracts*. https://office.clarin.eu/v/CE-2023-2328_CLARIN2023_ConferenceProceedings.pdf.
- Tomaž Erjavec, Maciej Ogrodniczuk, Petya Osenova, Nikola Ljubešić, Kiril Simov, Andrej Pančur, Michał Rudolf, Matyáš Kopp, Starkaður Barkarson, Steinþór Steingrímsson, Çağrı Çöltekin, Jesse de Does, Katrien Depuydt, Tommaso Agnoloni, Giulia Venturi, María Calzada Pérez, Luciana D. de Macedo, Costanza Navarretta, Giancarlo Luxardo, Matthew Coole, Paul Rayson, Vaidas Morkevičius, Tomas Krilavičius, Roberts Dargis, Orsolya Ring, Ruben van Heusden, Maarten Marx, and Darja Fišer. 2022. *The ParlaMint corpora of parliamentary proceedings. Language resources and evaluation*, 57:415–448.
- Norman Fairclough. 2013a. *Critical Discourse Analysis: The Critical Study of Language*. Longman applied linguistics. Taylor & Francis.
- Norman Fairclough. 2013b. *Language and Power*. Language In Social Life. Taylor & Francis.
- Darja Fišer and Jakob Lenardič. 2018. CLARIN Corpora for Parliamentary Discourse Research. In *Proceedings of the LREC2018 Workshop ParlaCLARIN: Creating and Using Parliamentary Corpora*. European Language Resources Association. http://lrec-conf.org/workshops/lrec2018/W2/summaries/14_W2.html.
- Goran Glavaš, Federico Nanni, and Simone Paolo Ponzetto. 2019. *Computational analysis of political texts: Bridging research efforts across communities*. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics: Tutorial Abstracts*, pages 18–23, Florence, Italy. Association for Computational Linguistics.
- Seth Jolly, Ryan Bakker, Liesbet Hooghe, Gary Marks, Jonathan Polk, Jan Rovny, Marco Steenbergen, and Milada Anna Vachudova. 2022. *Chapel Hill Expert Survey trend file, 1999–2019*. *Electoral Studies*, 75:102420.
- Johannes Kiesel, Çağrı Çöltekin, Maximilian Heinrich, Maik Fröbe, Milad Alshomary, Bertrand De Longueville, Tomaž Erjavec, Nicolas Handke, Matyáš Kopp, Nikola Ljubešić, Katja Meden, Nailia Mirzhakhmedova, Vaidas Morkevičius, Theresa Reitis-Münstermann, Mario Scharfbillig, Nicolas Stefanovitch, Henning Wachsmuth, Martin Potthast, and Benno Stein. 2024. *Overview of touché 2024: Argumentation systems*. In *European Conference on Information Retrieval*, pages 466–473. Springer.
- Taja Kuzman, Nikola Ljubešić, Tomaž Erjavec, Matyáš Kopp, Maciej Ogrodniczuk, Petya Osenova, Paul Rayson, John Vidler, Rodrigo Agerri, Manex Agirrezabal, Tommaso Agnoloni, José Aires, Monica Albini, Jon Alkorta, Iván Antiba-Cartazo, Ekain Arrieta, Mario Barcala, Daniel Bardanca, Starkaður Barkarson, Roberto Bartolini, Roberto Battistoni, Nuria Bel, Maria del Mar Bonet Ramos, María Calzada Pérez, Aida Cardoso, Çağrı Çöltekin, Matthew Coole, Roberts Dargis, Jesse de Does, Ruben de Libano, Griet Depoorter, Katrien Depuydt, Sascha Diwersy, Réka Dodé, Kike Fernandez, Elisa Fernández Rei, Francesca Frontini, Marcos Garcia, Noelia García Díaz, Pedro García Louzao, Maria Gavriilidou, Dimitris Gkoumas, Ilko Grigorov, Vladislava Grigorova, Dorte Haltrup Hansen, Mikel Iruskietia, Johan Jarlbrink, Kinga Jelencsik-Mátyus, Bart Jongejan, Neeme Kahusk, Martin Kimbauer, Anna Kryvenko,

Noémi Ligeti-Nagy, Giancarlo Luxardo, Carmen Magariños, Måns Magnusson, Carlo Marchetti, Maarten Marx, Katja Meden, Amália Mendes, Michal Mochtak, Martin Mölder, Simonetta Montemagni, Costanza Navarretta, Bartłomiej Nitoń, Fredrik Mohammadi Norén, Amanda Nwadukwe, Mihael Ojsteršek, Andrej Pančur, Vassilis Papavassiliou, Rui Pereira, María Pérez Lago, Stelios Piperidis, Hannes Pirker, Marilina Pisani, Henk van der Pol, Prokopis Prokopidis, Valeria Quochi, Xosé Luís Regueira, Michał Rudolf, Manuela Ruisi, Peter Rupnik, Daniel Schopper, Kiril Simov, Laura Sinikallio, Jure Skubic, Minna Tamper, Lars Magne Tungland, Jouni Tuominen, Ruben van Heusden, Zsófia Varga, Marta Vázquez Abuín, Giulia Venturi, Adrián Vidal Miguéns, Kadri Vider, Ainhoa Vivel Couso, Adina Ioana Vladu, Tanja Wissik, Väinö Yrjänäinen, Rodolfo Zevallos, and Darja Fišer. 2023. [Linguistically annotated multilingual comparable corpora of parliamentary debates in English ParlaMint-en.ana 4.0](#). Slovenian language resource repository CLARIN.SI.

Jakob Lenardič and Darja Fišer. 2023. CLARIN Resource Families: Parliamentary Corpora. <https://www.clarin.eu/resource-families/parliamentary-corpora>, accessed on 2024-01-20.

T.A. van Dijk. 2008. *Discourse and Power*. Bloomsbury Publishing.

Federico Vegetti and Daniela Širinić. 2019. Left-right categorization and perceptions of party ideologies. *Political Behavior*, 41(1):257–280.