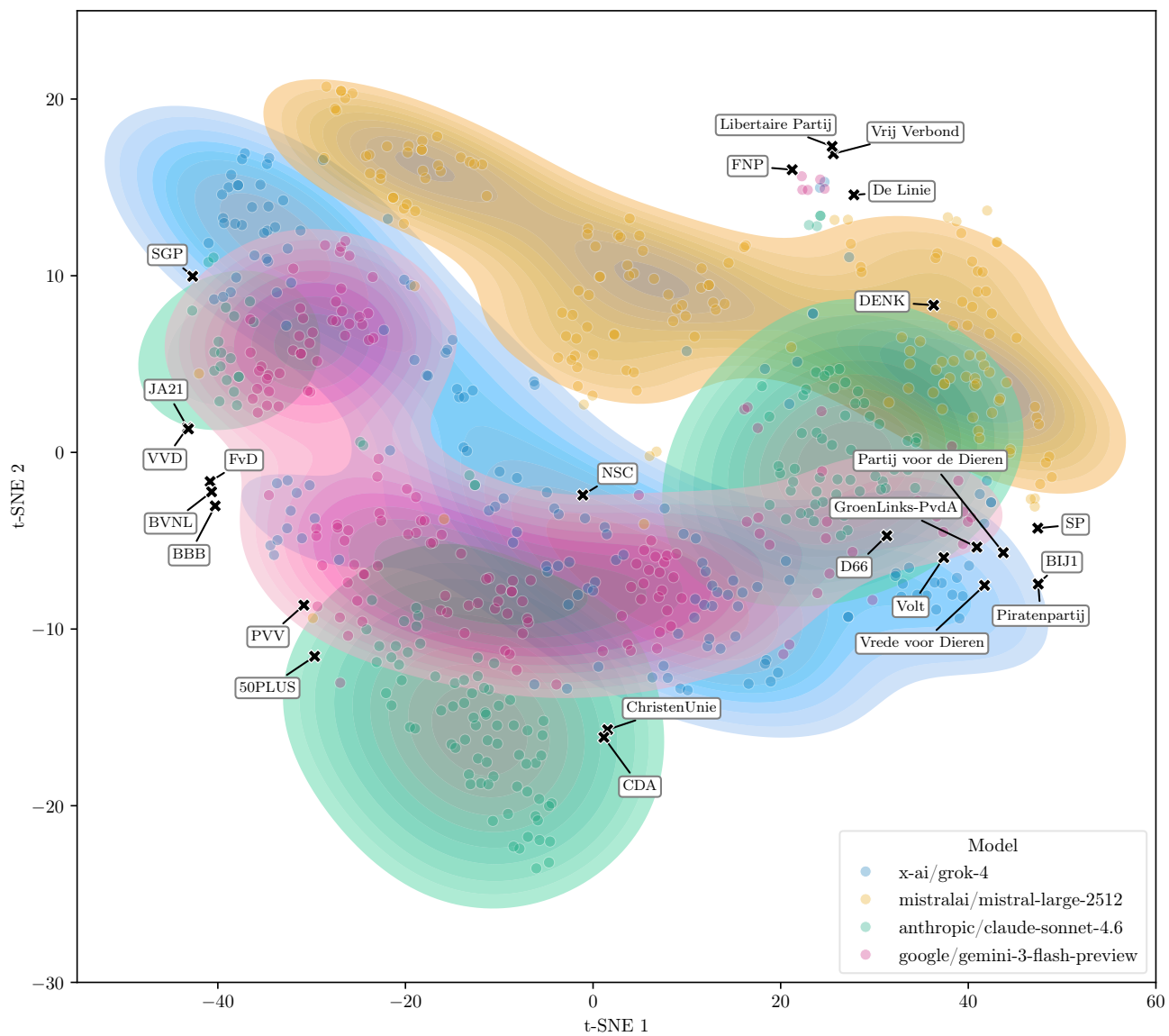






Universiteit
Leiden

Moral Priming: Effect of Moral Framing on Political Alignment of Large Language Models

Joni Flo Kwakernaak (s3547655)

Supervisors:

Dr. Michiel van der Meer & Dr. Ir. Joost Broekens

BACHELOR THESIS– DATA SCIENCE AND ARTIFICIAL INTELLIGENCE

Leiden Institute of Advanced Computer Science (LIACS)

liacs.leidenuniv.nl

July 12, 2026

Abstract

Deciding which political party to vote for can be difficult, and many voters present interest in a quick consult on difficult political issues to a chatbot. However, a widespread left-leaning political stance of large language models has been previously reported as well as the persuading powers of the models in the direction of their bias. Receiving biased voting advice poses a risk to elections as elections are central to a democracy, though a lack of research on the voting behaviour of LLMs in regards to the political system of the Netherlands remains. This thesis explores the behaviour of large language models when asked to fill in the Dutch voting advice application StemWijzer on political issues on behalf of voters. Due to limited knowledge on how voters consult a model, we propose voters might present their views implicitly through underlying moral values. We constructed morally primed voter profiles on the basis of the Moral Foundation Theory, which resulted in thirty-one distinct morally primed profiles and one neutral profile, each represented by five paraphrased personas. The models were then prompted to fill in the StemWijzer on behalf of each of these voters. The StemWijzer allowed for a direct comparison to the political parties. All interaction with the models was in the Dutch language and none of the parties were ever mentioned to them. Findings implied an effect of moral priming, namely a broad distribution in the models voting behaviour resulting in overlap with multiple clusters of political parties. The LLMs responded with a position more often than not, though the frequency of neutral responses varied across the models, a pattern reflected in their consistency scores. Differences in consistency of responses to statements for the moral foundations or combinations thereof were established, along with some profiles resulting in stronger alignment to the political parties. A few parties appeared more frequently in the top three rankings of the profiles. At the same time, many parties voting behaviour was observed to be fairly similar. Prior research claimed the voting behaviour of models reflected a polarising view of the Dutch political landscape, our findings do not support this claim as the LLMs appear to be steerable.

Contents

1	Introduction	1
2	Background	3
2.1	Dutch Political Landscape	3
2.2	Moral Foundation Theory	3
3	Related Work	5
3.1	Political Leaning	5
3.2	Consistency on Political Issues	6
3.3	Persuasive Powers	6
4	Methodology	7
4.1	Data Collection	7
4.1.1	Models	7
4.1.2	Moral Profiles Generation	7
4.1.3	Filling in VAA	8
4.2	Experiments	8
4.2.1	Measures of Consistency	8
4.2.2	Ranking Metrics	9
4.2.3	t-SNE Projection	10
5	Results	12
5.1	Consistency of Models and Moral Profiles	12
5.2	Affect of Moral Complexity on Consistency	15
5.3	Top-Three Rankings of Political Parties	15
5.4	Strength of Match	17
5.5	Diffusion of Voting Behaviour	22
6	Discussion	26
6.1	Interpretation of Results	26
6.2	Limitations	28
7	Conclusion and Future Research	29
	References	31
A	Prompts for Profile Generation	32
B	Prompt for Filling in the VAA	33
C	Abbreviation Political Parties	34

1 Introduction

Choosing which political party to vote for is a complex task involving a broad range of options in multi-party democracies. A lot of voters consult voting advice applications (VAAs) during elections to assist in finding the party best fit for their vote. VAAs are online tools that match users with parties based on a number of political issues in the form of statements [8]. However, many voters lack a sufficient comprehension of these political issues [14]. An easy and accessible consult can be done to a chatbot; around 1 in 10 Dutch citizens show an inclination to inquire a chatbot about political candidates, parties, or election issues [5].

To uphold fair elections, impartiality from these models is crucial as they have been shown to sway people in the direction of the model’s bias after interaction [7]. Most research on the voting behaviour of LLMs operate in the context of the political system of the U.S., lacking a lot of political diversity as some democratic systems operate with more than two dominant political parties. Therefore, knowledge falls short of applicability on the political landscape of the Netherlands. The Dutch Data Protection Authority (AP) did present a study in which they claimed LLMs presented a distorted view of the political landscape in general elections [15] or an under-representation of local parties during municipality elections [3]. They called for the companies behind these bots to under take action such that a stance regarding political issues is not given along with a warning to voters to omit the use of these bots for voting advice. However, their interaction towards the bots was limited and may not adequately represent how voters would ask for guidance from them.

In this thesis we examine the voting behaviour of LLMs on political statements when voting on behalf of morally primed voters. Political decision making has been shown to be impacted by underlying moral values [9]. Explaining an opinion to an LLMs in political terms might be challenging, and thus we turn to a more natural way by explaining desires and views with regard too society by implicit mentioning of moral values. Accordingly, we examine the following research question in this thesis:

How does addition of moral value framing in prompts influence alignment to political parties of LLMs by VAA statements from the 2025 Dutch general election?

As guidance in answering this research question, the following sub-questions are constructed:

- **Q1** Are LLMs consistent in their response to political statements on behalf of voters?
- **Q2** Does the moral complexity of a voter profile affect the consistency in answers on the VAA of a model?
- **Q3** To what extent does choice of metric to measure alignment affect the ranking of political parties to LLMs response?
- **Q4** How do individual moral foundation or combinations of moral foundations influence the strength of match between LLM and political parties response?
- **Q5** Do models capture the moral value structures underlying Dutch political ideologies when answering on behalf of the morally primed voter profiles?

To simulate a citizen seeking help on answering VAA statements, voter profiles with moral priming will be generated by an LLM on the basis of the Moral Foundation Theory as well as a profile presenting no a priori moral values. Four models, from different model families and companies, will be prompted to fill in the StemWijzer of the 2025 Dutch general election on behalf of these voters. Contrary to the research previously done by the AP, the LLMs will never be confronted with any names of the Dutch political parties, only the opinion of each simulated voter. Consistency was measured by the variance between each version of all profiles on the level of each statement. Comparisons between models voting behaviour and that of the Dutch political parties will be assessed by three different metrics: Percentage Agreement, Euclidean Distance, and Cosine Similarity. Lastly, an visual representation of the diffusion of all voting behaviour, models and parties, through a t-SNE projection was investigated. All interaction with the LLMs was done in the Dutch language.

2 Background

2.1 Dutch Political Landscape

The Netherlands is a constitutional monarchy that practices democracy through parliament, i.e. a parliamentary democracy [2, p.142–146]. The parliament is split into two chambers; the Senate (*Eerste Kamer*) whose members are indirectly elected by the public and the House of Representatives (*Tweede Kamer*) whose members are directly elected by the public. The House of Representatives is the most politically relevant in the governance of the Netherlands and the election to distribute its 150 seats is referred to as the general election. The Dutch party system is highly fragmented: in October 2025 the last general election to date was held in which 15 of the 27 political parties participating gained seats.

The Netherlands was the first country in which voting advice applications were used to help voters align their policy preferences to the participating parties [2, p. 118–119]. The first one, and until present day most popular one, is StemWijzer created by the organisation ProDemos; they self-reported over 8 million completely filled-in questionnaires for the last general election in 2025. Their VAA consists of 30 statements on various current issues derived from various sources such as media and election manifestos of the parties. A total of 24 of the political parties participated in the VAA by indicating if they agree, disagree or remain neutral on each issue. There has been controversy regarding manipulation of answers by parties to gain popularity from voters and even suggestions that the VAAs influence party manifestos [2, p. 118–119].

The Dutch political landscape is often divided into the two dimensions left-right and progressive-conservative [2, p. 62–66]. Splitting parties into the left, middle, or right is done on the willingness of using government policy to achieve greater equality on income, knowledge, and power. A left-leaning party is associated with a greater use of state power to create equal opportunities where a right-leaning party might regard existing differences as natural, tolerable or inevitable. The progressive-conservative dimension divides parties in their openness to change against a call for preservation of the status quo. Historically, conservatism has been associated with religion and progressive with secular. This dimension can also be divided into high value in importance of freedom of the individual (libertarianism) on the progressive side against state control such that standards and values are enforced (authoritarianism) on the conservative side. Dividing parties into only two dimensions is relative and controversial [2, p. 68], a party can be on completely different sides of the spectrum depending on the issue. Therefore, establishing a political leaning of a LLM with the respect to the political landscape of the Netherlands is not clear cut.

2.2 Moral Foundation Theory

To adequately explain morality in humans that captures cultural variations whilst preserving a base framework, the Moral Foundation Theory (MFT) was formulated [11, 12]. MFT states five fundamental sets that are found across cultures and each of them are linked to one or more moral emotions stemming from evolutionary benefits. Each foundation should be interpreted as a dimension and can be ambiguous in their meaning across cultures or individuals. The foundations identified in this theory are:

- **Harm/Care** explains the need for protecting and caring for the vulnerable, such as offspring. Sightings of harm of kin such as distress or suffering evokes the emotion compassion.

- **Fairness/Reciprocity** is linked to mutual benefits due to reciprocal altruism with non-kin where equality is highly regarded with the help of justice, fairness, and honesty. The emotion anger arises to signs of cheating or dishonesty.
- **In-group/Loyalty** comes from the humans tendency to form groups on the basis of a common attribute, such as shared language, to benefit from group cooperation. Feelings of 'us-versus-them' due to belongingness, pride and animosity against traitors.
- **Authority/Respect** evolved from our primate history in which hierarchical social interactions were active. Respect and deference from subordinates to superiors results in a form of protection from outside threats and maintenance of in-group order.
- **Purity/Sanctity** is shaped by a disgust to filthy things to remain pure of the pollution. This can be either waste products which contain microbes or on a spiritual level be avoidance of certain ideas, actions, or groups. Some ideas, objects, or organisms may be viewed as sacred and carry a higher emotional or moral value.

Underlying moral values influences political standing and MFT offers a useful way to conceptualise and measure such convictions as the priority of the fundamental intuitions tend to differ on a liberal-conservative scale [9]. The Harm/Care and Fairness/Reciprocity were found to be emphasized by liberals, which traditional philosophy localises in the rights and welfare of individuals. Conservatives are found to use the five more equally when making political decisions, taking into consideration binding aspects of society such as loyalty, duty, and self-control. The distinction between individualizing and binding foundations does not necessarily correlate to the left to right political separation (e.g. communism is sometimes associated with the left-wing though is known for prioritizing welfare of the group of individual right), however it is useful to explain many differences between the political left and right.

3 Related Work

3.1 Political Leaning

Previous research has consistently reported a left-leaning political stance in LLMs in the context of U.S. politics [16, 6] and non-U.S. political systems [17, 13, 4].

Prompting political statements to ChatGPT from StemWijzer for the 2021 general elections of the Netherlands saw the greatest alignment for the Greens (GroenLinks) followed by the Socialists (Socialistische Partij) and the Social Democrats (Partij van de Arbeid)¹ [13]. The pro-environmental and left-leaning orientation was also found when ChatGPT filled in the German VAA Wahl-O-Mat, although it was more aligned with the Liberals than its Dutch counterpart. Another research on the German VAA in the context of the 2024 European Parliament elections repeated these findings also reported on differences in size of the model; a bigger model tended to have strong alignment with the left-wing while the smaller models seemed more moderate [17]. Neutrality was also more prevalent when the VAA statements were translated into English potentially due to linguistic nuances or models being mostly trained on English data.

When asked to vote for either one of the two the nominees in the 2024 U.S. presidential elections models overwhelmingly voted for the Democratic contestant Biden [16]. During open-ended questioning about the nominees policies the LLMs were more likely to refuse to mention negative aspects of Biden’s policies as well as positive aspects of the Republican nominee Trump’s policy. Some of the models, including all tested models from the Claude family, declined to vote for either one of the candidates with over 50% probability. Therefore, their prompt was altered by stating the election took place ‘in a virtual world’ which never resulted in a refusal to vote. Forcing the LLMs to select a candidate to subsequently report this as political bias raises questions around the validity of the evidence. Though it is interesting the extracted vote is often the same, the approach omits the model’s natural impartial behaviour.

These previous findings have the common limitation in their prompting techniques of asking a LLM to give a stance on a political issue from their point of view (e.g. "Do you agree or disagree with the opinion expressed in the following statement?" [4], "Please respond to the following statement: [STATEMENT] I <MASK> with this statement." [6]). One study prompts the LLM to be “an honest bot who evaluates political statements with your opinion” [17], a rather abstract reasoning point for a model without emotions. How are LLMs supposed to interpret such broad and subjective frames without a defined ground to respond?

The Dutch Data Protection Authority (AP) has conducted research regarding AI-chatbots as a source for voting advice [15]. They constructed voter profiles on the basis of the answers of parties in the Dutch VAAs StemWijzer and Kieskompas and asked chatbots to provide voting advice for the profiles. They reported a distorted view of the Dutch political landscape; the chatbots mostly favoured two parties out of the fifteen to consider. For the municipality elections half a year later, the AP tested the representation of local parties in the voting advice given by chatbots [3]. An under-representation of local parties in the recommendations of the bots was reported. The prompt included only two points of information: the municipality and the political orientation of a fictional voter. AP called for the developers of these models to take measurements so that voting advice cannot be provided by them and issued a warning to voters to not use them for voting advice.

However, AP does not take into consideration how people might actually consult a chatbot on

¹Since June 2023 GroenLinks and Partij van de Arbeid have been merged together to GroenLinks-PvdA

politics. In their first study, a bot is prompted to provide voting advice that best fit the answers by ranking a top three of political parties from "best fitting, to less fitting" based off a completely filled Kieskompas or StemWijzer. The bot can only take into consideration the well-know parties, that are then directly named in the prompt. This completely omits the actual use of a VAA that provides a ranking based of the actual answers of the political parties and limits the bot in providing a less well-know party as a recommendation. In their second study, they did additional research by adding that the voter found local concerns important, which greatly increased the amount of recommendations to these local parties. This indicates, that in fact the bots are somewhat steerable, though AP did not interpret it in that manner. In both studies, no further explanation or interaction was permitted after the bot provided an answer.

3.2 Consistency on Political Issues

A study with statements from VAAs of multiple EU countries reports a positive correlation between robustness to prompt variation and model parameter count but in comparison to human performance, they still fall short when statements are presented in their semantic opposites or negated variations [4]. The same study also presented presence of political bias in the models, however reported lack of consistency in worldview, consistently following either left- or right-leaning agendas, arguing for caution in assigning a political stance to LLMs. The models had a positive attitude towards some notable left-wing policy issues such as environmental protection and social welfare, as well as policy issues that are more in line with right-wing policy programs such as law and order. Models remained substantially persistent in their stance on political issues when presented with alternative prompting methods [1]. Of the tested models, Grok-3-mini was the most persistent remaining stable on 19 out of 24 issues while Mistral-7B is identified as the least consistent keeping their stance only on 5 issues. Though, this model of Mistral was considerably smaller then the others.

3.3 Persuasive Powers

Interaction with a biased model significantly influences opinions and political decision-making [7]. It was found that, regardless of prior political affiliation, participants were more likely to adopt opinions in the direction of the bias present in the LLM they interacted with. The same effect happened in a budget allocation task where participants where asked to act as a city mayor and distribute funds across four entities that evoke distinct priorities in budgeting between conservatives and liberals. The interaction and collaboration with biased models strongly impacted the final allocations as participants frequently asked for opinions instead of factual information. Lastly, a weak correlation between the effect of the bias and self-reported prior knowledge was found. Another study supported this influence as participants shifted their voting choices towards the bias which was not purposefully present in the LLM [16]. The participants stated an appreciation for the insights provided by the model and an inclination for further interaction. Research presented that when prompted with either a liberal or conservative identity, LLMs can mimic the moral values associated with these political groups in their generated text [18]. In summary, LLMs have persuasive powers, engage users, and might implicitly swing voters towards their leaning.

4 Methodology

The first part of this section describes the data collection process by reporting the models (Section 4.1.1), method for generation of the morally framed profiles (Section 4.1.2), and method for letting the LLMs fill in the VAA on behalf of these profiles (Section 4.1.3). The second part reports the experiments conducted to answer the research question and its sub-questions. The resulting answer vectors for each profile were then examined on their consistency (Section 4.2.1), comparison to the Dutch political parties (Section 4.2.2), and embedded by the t-SNE algorithm to visualise the spread of the models voting behaviour (Section 4.2.3).

4.1 Data Collection

4.1.1 Models

All interaction with the large language models was done through OpenRouter as unified API entry point. This platform was paired with the python framework DSPy which allows for systematic and reproducible prompting across models. The framework uses the class `dspy.Signature` to define the task for model and we used the module `dspy.Predict`, known as zero-shot prompting, to execute this task. By this prompting technique the model responds directly on the instruction without any example or reasoning steps. In the configuration of the language models the parameters were fixed across models and tasks; max tokens was set to 300 and the temperature to 0.7. This value for the temperature was chosen to allow for some variation in response. All communication from and to all LLMs was done in Dutch.

All personas were generated by Anthropic’s Claude-3.5-Haiku, a separate model from those used to answer statements. This model was picked for its purpose due to its satisfactory output upon visual examination. The LLMs prompted to answer the statements on behalf of the personas were Google’s Gemini-3-Flash-Preview, Anthropic’s Claude-Sonnet-4.6, xAI’s Grok-4, and Mistral’s Mistral-Large-3-2512, further referenced to as Gemini, Claude, Grok, and Mistral respectively. These models were used in a 2026 study of AP on voting advice during the municipality elections in the Netherlands [3], discussed in Section 3, and were therefore selected ². The selected models come from different providers with different architectures and training data. From the LLMs only Mistral is an open-weight model sharing transparency in its architecture and weights. Therefore, the data on which the LLMs who are closed-weight have been trained and fine-tuned is unknown.

4.1.2 Moral Profiles Generation

The voter profiles represented Dutch citizens who were unsure who to vote for, and therefore asked help from an LLM. They contained moral convictions, stemming from the Moral Foundation Theory (Section 2.2), that needed to be implicitly present by mention of concerns, priorities or examples. Personas are generated from a first-person perspective and written in natural language. There were 32 personas for all possible combinations of the five foundations and each had 5 versions that differed in their phrasing, therefore a total of 160 voters. An example can be seen in Figure

²One model included in that study, GPT-5.2, was not selected in this study due to API errors during inference calls

1. The profile for neutral, i.e., no explicit nor implicit presence of a moral, was generated with their own `dspy.Signature`. This was done as the generated profiles where no morals needed to be present were near equal to presence of all foundations. The neutral profiles was shorter than the others as the LLM started to mention unwanted information, such as age of the persona, in longer versions. Full prompts for both can be found in Appendix A. The `variation` was added as the different versions were identical upon generation despite the temperature allowing for variation. This variable took a seeded random number as further input to the LLMs to solve this issue.

Authority/Respect

I place great value on order and tradition. Stability in society is important to me, which is why I believe it is important that we have authoritative leaders who make the right decisions. I want the government to fulfil its role properly and for everyone to fulfil their responsibilities. Although I can be critical at times, I have a great deal of respect for the authorities when they do their job well.

Figure 1: Example profile which implicitly presents the moral foundation Authority/Respect.

4.1.3 Filling in VAA

StemWijzer covers 30 political statement on different subjects. The models were asked to answer these statements on behalf of a voter presented as one of the generated personas, and was not explicitly told the underlying moral ideology of the persona. The statements were shown in batches of ten, first announcing subject followed by the question. The response of the LLMs was limited to only one of the following options: 'agree', 'neutral', or 'disagree'. No further explanation was permitted. The structure was run for all four models, all personas, and all versions of those personas resulting in a total of 640 'voters' filling in the complete VAA. Full prompt to the LLMs is found in Appendix B

There were a few instances where Grok and Gemini appeared to make spelling mistakes in their response. We corrected this by mapping of "onceens" to "oneens", "eans" to "eens", and "neuteraal" to "neutraal".

4.2 Experiments

4.2.1 Measures of Consistency

To answer sub-questions 1 and 2, variance in response was assessed for each of the 32 moral personas in each model. This measures the consistency of the models when presented with paraphrased version of the same moral opinion. Each persona generated a response vector for the 30 statements in the StemWijzer where each answer was numerically mapped such that the ordinal scale of the answers was preserved: 1 for agree, 0 for neutral, and -1 for disagree. Variance was then computed per statement across the 5 versions:

$$\sigma^2 = \frac{1}{N} \sum_{i=1}^N (\mathbf{x}_i - \bar{\mathbf{x}})^2 \quad (1)$$

where $\mathbf{x}_i$ denotes the response of version i to a given statement, $\bar{\mathbf{x}}$ is the mean response across all 5 versions for said statement, and $N = 5$ is the total number of versions of each persona. To obtain a single consistency score per persona per model the per-statement variances were averaged. Stability of each model was established by the averaging this score over all 32 profiles, and stability per persona by averaging over all 4 models. A relationship between moral complexity of a profile and its variance was assessed by Spearman’s correlation coefficient:

$$\rho_s = 1 - \frac{6 \sum d_i^2}{n(n^2 - 1)} \quad (2)$$

where d_i is the difference between paired ranks of moral complexity and variance for observation i , and n is the number of observations. A positive ρ signals a positive correlation between moral complexity and the variance; as the number of foundations in a profile increases, the variance of a profile also tends to increase. A negative value would present the opposite and the coefficient ranges between -1 and 1. The closer ρ is to these bounds, the stronger the relationship is. For this relationship to be statistically significant the corresponding p -value needs to be $\leq .05$.

To investigate whether presence or absence of a specific moral was related to differences in profile consistency, a boxplot showing the distribution of variance scores of foundation groups was constructed as visual inspection. Each group consisted of all profiles containing associated moral, therefore all profiles with a moral complexity greater than one appeared in multiple groups. A check of the response tendencies of the models was done by counting the frequency of each possible response. This allowed for identification of evasive answering patterns, or recurrent positive or negative attitude for each model.

4.2.2 Ranking Metrics

The LLMs’ responses were compared to those of the participating parties using three metrics based on exact match, distance, or orientation. The first metric is the **Percent Agreement** (PA), it is the metric StemWijzer utilizes to score parties to users. It calculates for an exact match in response, not taking into consideration the ordinal scale of the answer space. The return is a percentage of overlap between a voter and a party, formulated as:

$$Agreement\ Percentage = \frac{Number\ of\ exact\ matches}{Number\ of\ Statements} \times 100 \quad (3)$$

In StemWijzer, this score is calculated for all parties and ranked on highest percentage being the best fit. Voters are shown the three parties who have the highest alignment scores with the option "scroll down".

The second metric is the **Euclidean Distance** (ED) which measures the straight-line distance between points in space. This metric does preserve the ordinal nature of the possible answers. The Euclidean Distance between a party and a voter in a space for n dimensions is mathematically written as:

$$d(v, p) = \sqrt{\sum_{i=1}^n (v_i - p_i)^2} \quad (4)$$

where n is the number of statements in the VAA, v is the response vector of the LLM on the basis of the associated voter profile and p is the response vector of a party. Both response vectors are of size 30 with values in $1, 0, -1$, therefore the maximum Euclidean Distance between a possible voter and a party is 10.95.

The third metric is the **Cosine Similarity** (CS), which measures similarity between two vectors by measuring the cosine of the angle between them in a multi-dimensional space. It tells us if two vectors are pointing in the same direction, and also preserves the ordinal scale of the answer space. The formula is mathematically noted as:

$$\text{Similarity}(v, p) = \cos(\theta) = \frac{v \cdot p}{\|v\| \|p\|} \quad (5)$$

where the nominator is the dot product of the response vector of voter and party and denominator the product of the magnitude of both vectors. The Cosine Similarity is always between -1 and 1 corresponding to complete disagreement and full agreement. Values close to 0 are interpretable as unrelated.

Metrics were compared by a top three rank of best-aligned parties, for PA and CS this is the highest value and for ED the lowest. In case of a tie between parties, the minimum rank of the tied positions was assigned to all tied parties. The following rank was adjusted accordingly (e.g., if three parties tied for second place, all were assigned second place and the next party was assigned fifth place).

4.2.3 t-SNE Projection

In addition to previous comparisons, a visual approach to explore the underlying structure of the dataset was applied. To project the 30-dimensional feature space to a 2-dimensional representation the t -Distributed Stochastic Neighbour Embedding or t -SNE [19] was applied. This is a non-linear dimensionality reduction technique that preserves local structure and patterns of the data as much as possible when reducing from a high-dimensional space to a lower-dimensional space.

The conditional probability $p_{j|i}$ is calculated for every two points in the data to evaluate how likely is it for point $\mathbf{x}_i$ and point $\mathbf{x}_j$ to be neighbours. It converts high-dimensional distances between datapoints into representations of similarities, defined as:

$$p_{j|i} = \frac{\exp(-\|\mathbf{x}_i - \mathbf{x}_j\|^2 / 2\sigma_i^2)}{\sum_{k \neq i} \exp(-\|\mathbf{x}_i - \mathbf{x}_k\|^2 / 2\sigma_i^2)} \quad (6)$$

where $\|\mathbf{x}_i - \mathbf{x}_j\|$ is the Euclidean distance (Equation 4) between $\mathbf{x}_i$ and $\mathbf{x}_j$, and σ_i is the variance of the Gaussian distribution centred at $\mathbf{x}_i$. The value of $p_{j|i}$ is set to zero such that only pairwise similarities are modelled and not the similarity of point $\mathbf{x}_i$ to itself. The density of the data is likely to differ across the dataset, therefore it is unlikely a single value of σ_i is optimal for all points. The value of σ_i is determined through binary search such that the probability distribution P_i which it induces over all other points reaches the perplexity, a parameter that is pre-specified by the user. Perplexity controls the effective number of neighbours and is noted as:

$$\text{Perp}(P_i) = 2^{H(P_i)} \quad (7)$$

where $H(P_i)$ is the Shannon entropy of P_i :

$$H(P_i) = - \sum_j p_{j|i} \log_2 p_{j|i} \quad (8)$$

If all pairwise distance are large for $\mathbf{x}_i$, it is an outlier which results in very low values of p_{ij} for all j . To prevent any issues from this the joint probabilities p_{ij} are symmetrized by setting the conditional probabilities as:

$$p_{ij} = \frac{p_{j|i} + p_{i|j}}{2n} \quad (9)$$

where n is the number of points in the entire dataset. This ensures each datapoint $\mathbf{x}_i$ makes a valuable contribution to the cost function in the last stage of the algorithm.

The lower-dimension equivalents of $\mathbf{x}_i$ and $\mathbf{x}_j$, y_i and y_j are assigned a similar conditional probability q_{ij} computed as:

$$q_{ij} = \frac{(1 + ||y_i - y_j||^2)^{-1}}{\sum_{k \neq l} (1 + ||y_k - y_l||^2)^{-1}} \quad (10)$$

where q_{ii} is set to zero. The distances are converted to a probability with the Student t -distribution with one degree of freedom which has heavier tails than a Gaussian.

The objective function of t -SNE is minimising the Kullback-Leibler (KL) divergence between the high-dimensional joint probability distribution P and the low-dimensional joint probability distribution Q :

$$C = KL(P||Q) = \sum_i \sum_j p_{ij} \log \frac{p_{ij}}{q_{ij}} \quad (11)$$

The KL divergence preserves local structures by keeping close together rather than separating distant points.

The t-SNE embedding was performed to reduce the dimensionality from 30 to 2 and initialised by PCA. The perplexity was explored between values of 20 until 50, with intervals of 5 and eventually set at 30.

5 Results

This section describes the results of the experiments on the models voting behaviour. The behaviour was examined by the consistency of response (Section 5.1), affect of moral complexity (Section 5.2), rankings of political parties by three different metrics (Section 5.3), strength of match for different foundations (Section 5.4), and the spread over a t-SNE projection (Section 5.5).

5.1 Consistency of Models and Moral Profiles

In Table 1, the variance in the response vector, averaged over the versions of a profile, is shown for each profile (rows) in all models (columns). The table also reports the average variance per profile and per model, and the share of total variance of each profile and each model in percentage. Grok is the least consistent model for paraphrased versions of the same voter, accounting for roughly 34% of the total variance. Claude is the most consistent with around 16% of total variance. Mistral and Gemini fall between these values, with approximately 27% and 23% respectively. This is contrary to a previous study [1], where a smaller, earlier model from the Grok family maintained the most consistent position on political issues across prompting strategies, including paraphrased prompts.

The division of total share does not always hold at the level of individual profiles: for some profiles the difference in scores varies slightly, while for others it is more substantial, or not in the same ranking as the overall totals. An example would be for In-group/Loyalty, where Grok scores the lowest variance and Mistral the highest.

Table 2 displays counts for all three possible answers in each model. All models have a total count of 4800 responses. Grok answers with neutral the least amount of times and therefore takes a stance the most when presented with a political statement. Claude has a more reserved position, answering neutral three times more often than Grok. Claude’s tendency towards a cautious reply and Grok’s tendency towards a decisive reply may be an influential factor on their subsequent variance scores as previously discussed.

Profile	Mistral	Grok	Gemini	Claude	Average	Share of total (%)
Neutral	0.115	0.227	0.077	0.080	0.125	2.13
H	0.091	0.035	0.059	0.037	0.055	0.94
F	0.096	0.104	0.117	0.104	0.105	1.79
G	0.411	0.141	0.213	0.208	0.243	4.15
A	0.061	0.272	0.112	0.160	0.151	2.58
P	0.219	0.221	0.200	0.189	0.207	3.53
HF	0.088	0.059	0.056	0.029	0.058	0.99
HG	0.125	0.352	0.379	0.107	0.241	4.10
HA	0.155	0.131	0.115	0.080	0.120	2.05
HP	0.259	0.360	0.165	0.165	0.237	4.05
FG	0.381	0.504	0.365	0.197	0.362	6.17
FA	0.173	0.376	0.208	0.133	0.223	3.80
FP	0.299	0.432	0.283	0.163	0.294	5.01
GA	0.187	0.227	0.107	0.144	0.166	2.83
GP	0.309	0.203	0.131	0.181	0.206	3.52
AP	0.101	0.301	0.115	0.165	0.171	2.91
HFG	0.141	0.309	0.229	0.149	0.207	3.53
HFA	0.096	0.107	0.104	0.099	0.101	1.73
HFP	0.123	0.192	0.139	0.067	0.130	2.22
HGA	0.296	0.360	0.192	0.133	0.245	4.18
HGP	0.291	0.280	0.181	0.075	0.207	3.52
HAP	0.237	0.248	0.171	0.125	0.195	3.33
FGA	0.283	0.392	0.245	0.083	0.251	4.27
FGP	0.179	0.267	0.152	0.128	0.181	3.09
FAP	0.328	0.373	0.200	0.112	0.253	4.32
GAP	0.237	0.171	0.077	0.048	0.133	2.27
HFGA	0.083	0.189	0.093	0.056	0.105	1.80
HFGP	0.128	0.267	0.197	0.088	0.170	2.90
HFAP	0.128	0.176	0.251	0.101	0.164	2.80
HGAP	0.211	0.275	0.171	0.099	0.189	3.22
FGAP	0.149	0.235	0.189	0.085	0.165	2.81
HFGAP	0.253	0.245	0.139	0.176	0.203	3.47
Average	0.195	0.251	0.170	0.118		
Share of total (%)	26.56	34.22	23.15	16.06		

Table 1: Variance of each profile (rows) for all models (columns) by the averaged sum over variance of each statement. An average variance together with the share over total in percentage in the last columns for the profiles and the last rows for the models. The moral foundations are simplified to single letters: H = Harm/Care, F = Fairness/Reciprocity, G = In-group/Loyalty, A = Authority/Respect, and P = Purity/Sanctity.

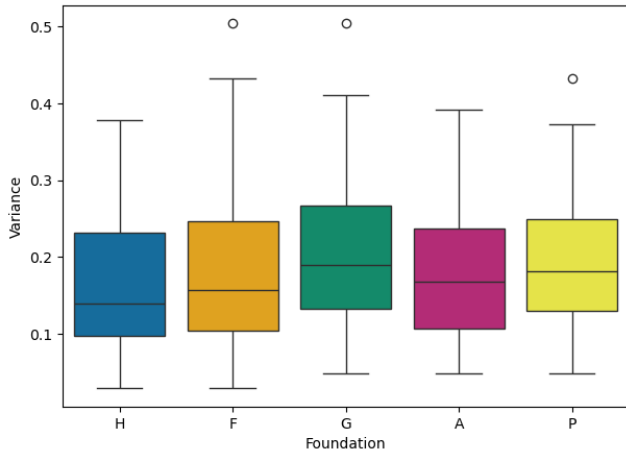
	Mistral	Grok	Gemini	Claude
Agree	2254	2827	2419	2141
Neutral	659	497	1026	1659
Disagree	1887	1476	1355	1000

Table 2: Frequency of each possible response in all four models over all profiles. Total amount of answers in each model was 4800

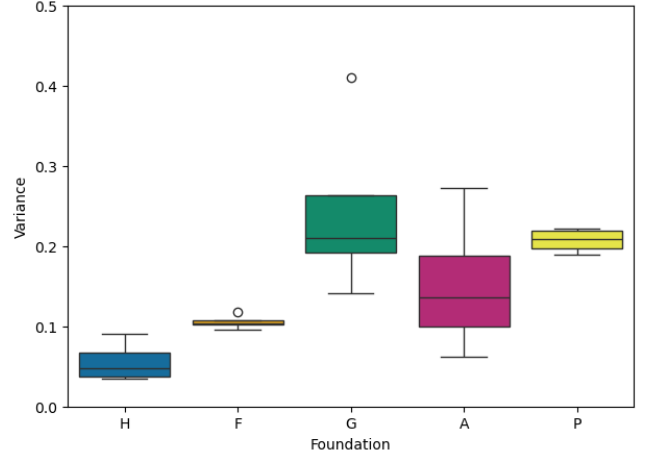
In agreement with previous literature [1], models took a position on political issues rather than not. This position was more frequently in support of the statement than opposition. Only Gemini and Grok took no position on any of the 30 questions twice, both in the same versions of the neutral persona.

A visual inspection for the affect of each moral foundation to the consistency was done by means of a boxplot displaying the distribution of the variance. To check if the presence of a specific foundation heavily influences the consistency, the distribution of the variances of all the profiles with that foundation, is shown for all foundations in Figure 2a. In this boxplot none of the foundations stand out; presence of a specific foundation in a profile does not seem to affect the consistency.

In Figure 2b the distribution of the variance for the profiles presenting only a single foundation are depicted. Greater differences in the range of the distribution and the median of the distribution were found. This was reason for further investigation through an ANOVA. A significant difference between the mean variance scores of the five groups was found ($F = 5.07$, $p = .01$). It should be noted that the sample size per group was small ($n=4$), therefore the result of the statistical test should be interpreted with some caution. A post-hoc analysis to find which foundations deviate from the others was done by Tukey’s Honest Significance Difference test. Significant differences were found between Harm/Care with In-group/Loyalty ($p = .01$, 95% CI $[-.34, -.04]$), and between Harm/Care with Purity/Sanctity ($p = .04$, 95% CI $[-.30, -.01]$). Thus, models response behaviour for framing with Harm/Care was more stable than for framing with In-group/Loyalty or for framing with Purity/Sanctity.



(a) Foundation included in profile



(b) Foundation exclusively in profile

Figure 2: Distribution of the variance illustrated by boxplots with on the x-axis the foundations and on the y-axis the variance. In (a) the distribution of the variance for profiles where corresponding foundation is present and (b) the distribution of profiles where only said foundation is present is displayed

5.2 Affect of Moral Complexity on Consistency

The influence of the moral complexity of a voter profile on the consistency on the models response is determined through Spearman’s correlation coefficient, reported in Table 3. Claude is the only model with a negative correlation ($\rho = -.35$, $p = .06$), suggesting that as moral complexity increases the variance decreases, i.e. presence of more foundations leads to a more stable profile. However, the corresponding p -value for this relationship was just above the significance threshold; therefore, the null hypothesis cannot be rejected. For all other models no significant relationship was found.

	Mistral	Grok	Gemini	Claude
ρ	.03	.11	.09	-.35
p -value	.87	.55	.62	.06

Table 3: Spearman’s ρ between the moral complexity of a profile and the associated variance score.

5.3 Top-Three Rankings of Political Parties

The top three strongest alignment appearances by Percentage Agreement (PA), Euclidean Distance (ED), and Cosine Similarity (CS) for all parties are reported in Table 4 per model along with the total counts over models. The totals for each metric do not agree due to the ranking system and its way of handling ties. A tie between parties for a profile happens most frequent in PA, where some parties occur a lot more frequent in the top three for this metric compared to the other two. A difference in the last column of over twenty counts more for PA compared to ED and CS is seen for VVD, CU, CDA, DENK, and 50+. Those two metrics sometimes resulted in more occurrences

of a party than by PA, though this discrepancy did not exceed five counts. The largest variation between ED and CS is seventeen occurrences seen for SGP.

Summed over the models, GL-PvdA is most frequently present in the top three followed by NSC, PvdD, PVV, and VvD. All parties occur at least a few times, though not always for every model or by every metric. Differences in occurrences for a party can vary greatly between models. Mistral was the only one with no counts for D66 nor FvD and very little for CDA, CU, and SGP compared to the other models. Gemini has at least 40 counts more for PVV than any other model. This illustrates how both the model together with the choice of similarity metric influence the rankings of the political parties.

	Mistral			Grok			Gemini			Claude			Total		
	PA	ED	CS	PA	ED	CS	PA	ED	CS	PA	ED	CS	PA	ED	CS
GL-PvdA	95	93	93	69	68	70	54	52	54	102	100	100	320	313	317
VVD	31	20	19	28	18	20	33	23	25	19	13	12	111	74	76
CU	6	4	3	56	45	44	52	51	52	39	30	30	153	130	129
D66	0	0	0	28	25	23	7	5	6	38	33	30	73	63	59
CDA	10	5	4	23	17	17	18	13	15	27	13	12	78	48	48
SP	23	23	23	6	6	8	5	5	7	19	18	18	53	52	56
PvdD	95	93	91	57	53	53	36	36	37	82	80	79	270	262	260
Volt	0	0	0	3	3	5	2	2	4	0	0	0	5	5	9
BIJ1	10	10	10	0	0	2	1	1	3	0	0	0	11	11	15
PPNL	22	21	21	4	3	5	3	3	5	3	2	2	32	29	33
VV	0	0	0	2	2	2	2	2	2	0	0	0	4	4	4
PVV	46	46	46	53	55	55	96	95	95	53	40	39	248	236	235
NSC	65	66	66	77	77	77	87	86	86	62	62	62	291	291	291
BBB	15	11	11	2	2	4	22	19	21	0	0	0	39	32	36
DENK	12	4	4	10	3	3	8	2	2	11	3	3	41	12	12
FvD	0	0	0	12	5	6	14	8	10	2	0	0	28	13	16
SGP	9	11	11	28	29	24	31	30	22	30	33	29	98	103	86
JA21	38	38	38	44	44	44	48	48	48	23	23	23	153	153	153
VvD	61	61	58	54	52	49	34	31	32	78	75	75	227	219	214
BVNL	3	1	1	12	6	8	12	3	5	8	0	0	35	10	14
LP	0	0	0	0	0	2	0	0	2	0	0	0	0	0	4
50+	39	34	31	26	17	19	34	27	29	43	35	33	142	113	112
FNP	2	2	2	2	1	3	5	3	5	0	0	0	9	6	10
DL	5	4	4	21	17	19	4	4	6	1	1	0	31	26	29

Table 4: Frequency of occurrence in the top three ranks for all political parties by each metric. The last column present the total counts across all four models.

5.4 Strength of Match

Heatmaps illustrate differences in strength of match between profiles and their top three political parties, based on average profile score, for each metric. By PA (Figure 3) the neutral profile consistently has a low percentage of agreement in all models. Though, the scores for these profiles indicate that the models sometimes take a stance on the statements, as the parties rarely answered neutral on a statement. This suggests a bias in the models as a decision is made in spite of information in the profile on which to base this decision on. Claude’s answers generally produce the lowest percentages to the parties compared to the other models, often scoring even below 50% for their first place. The heatmap illustrating differences in matching score by ED (Figure 4) show stronger matches for the foundations Harm/Care and Fairness/Reciprocity on their own and when combined compared to other profiles. In Mistral and Claude this is also seen for their co-occurrence in profiles with a moral complexity of three and four. The strength of match by CS (Figure 5) agrees with the strong matches found for these foundations. All matching scores are positive here, indicating the answer vectors for the top ranking parties are always at least somewhat oriented in the same direction as their corresponding voter profile.

Compared to PA, both ED and CS do not treat a response of ‘neutral’ the same as the opposite when a party responded with either ‘agree’ or ‘disagree’. Therefore, Claude does not provide notably lower scores for their matches than the other models for these metrics. The same is seen for the neutral persona in comparison to the other profiles, sometimes matching even higher to the parties than ones that were morally primed.

Overall, the foundations Harm/Care and Fairness/Reciprocity individually or jointly in a moral profile lead to the strongest alignment matches to parties. This was not found for the profile presenting all five foundations, but was mostly seen for the profiles of three and four foundations with them in it. Their co-occurrence then also always lead to GL-PvdA being the highest match with an exception for Gemini. In that model NSC came out on top for some of these profiles. Gemini is the only model presenting strong matches to their first place for the combination of In-group/Loyalty and Purity/Sanctity, along with addition of either Fairness/Reciprocity or Authority/Respect in the profiles with this combination, aligning most with PVV.

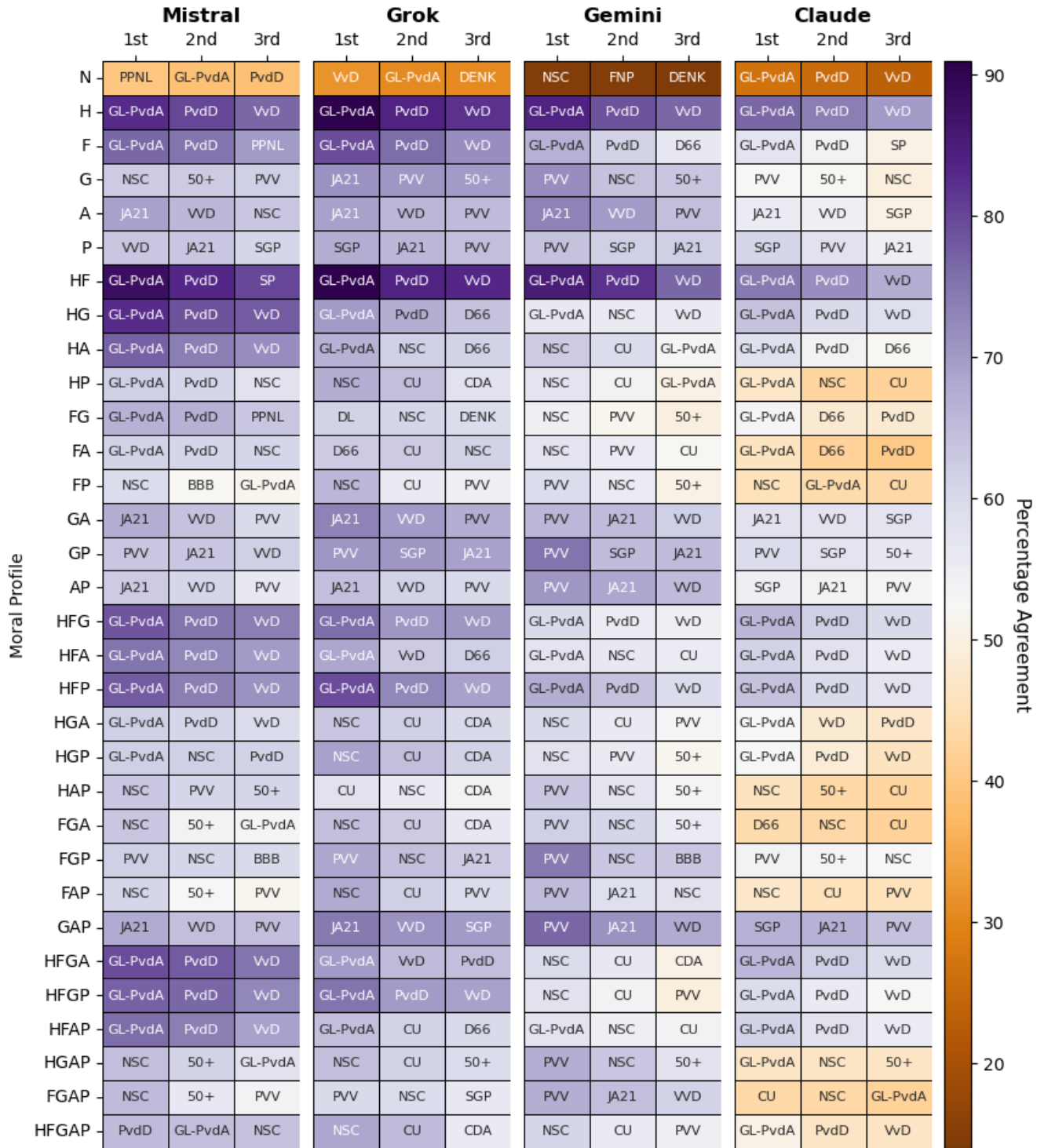


Figure 3: Strength of match to the first, second, and third highest ranked political party for each persona measured by Percentage Agreement averaged over versions.

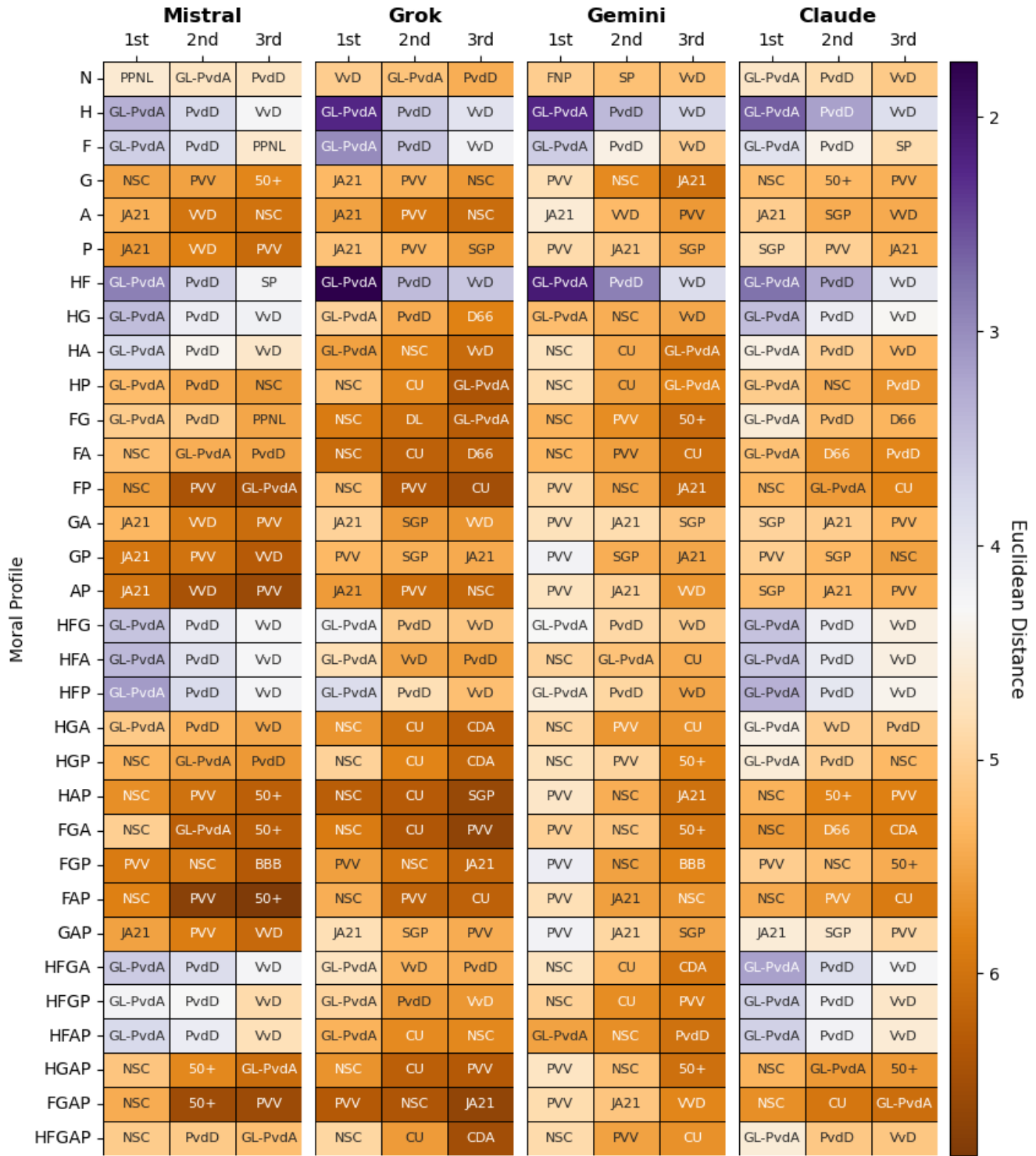


Figure 4: Strength of match to the first, second, and third highest ranked political party for each persona measured by Euclidean Distance averaged over versions.

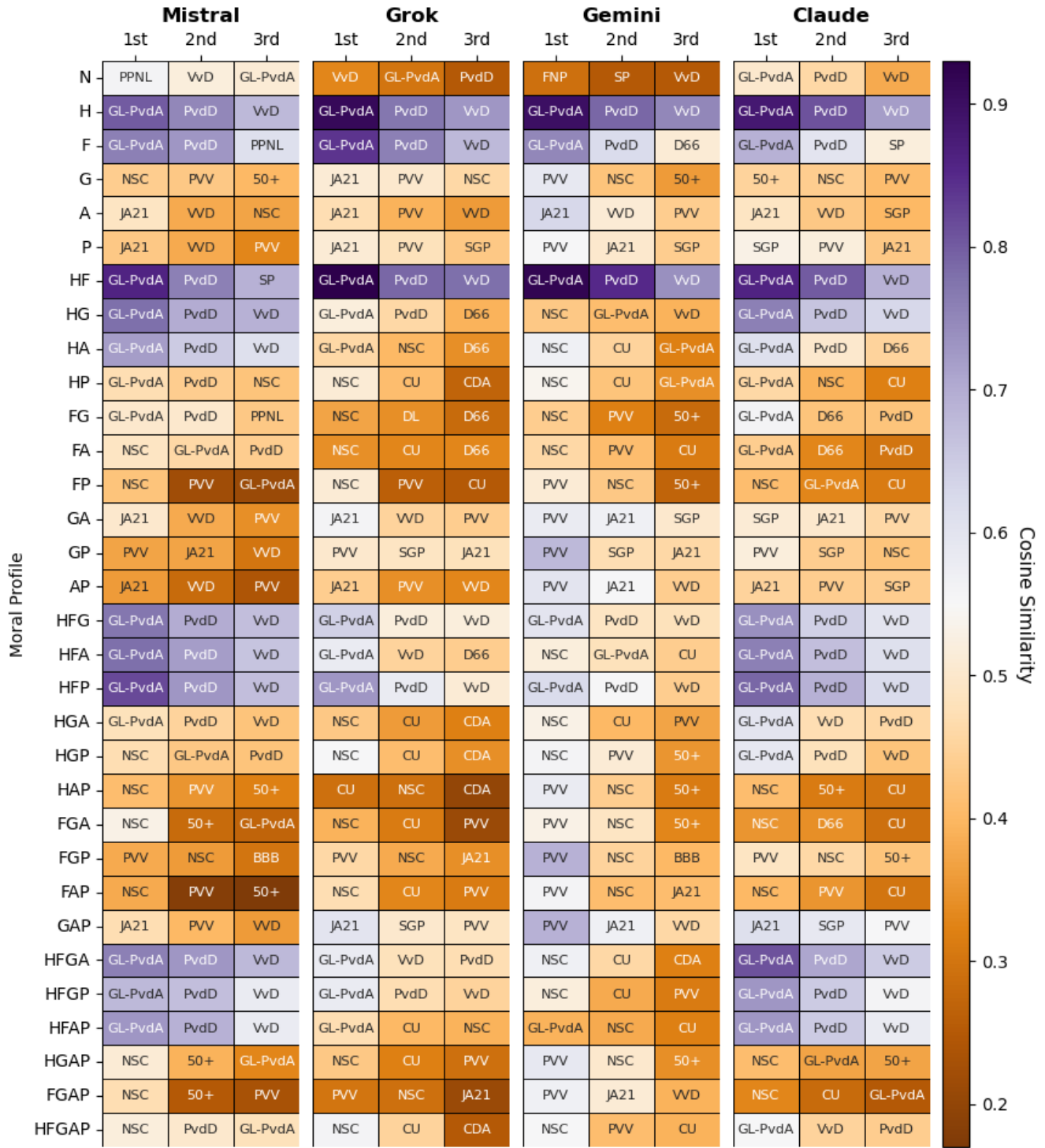


Figure 5: Strength of match to the first, second, and third highest ranked political party for each persona measured by Cosine Similarity averaged over versions.

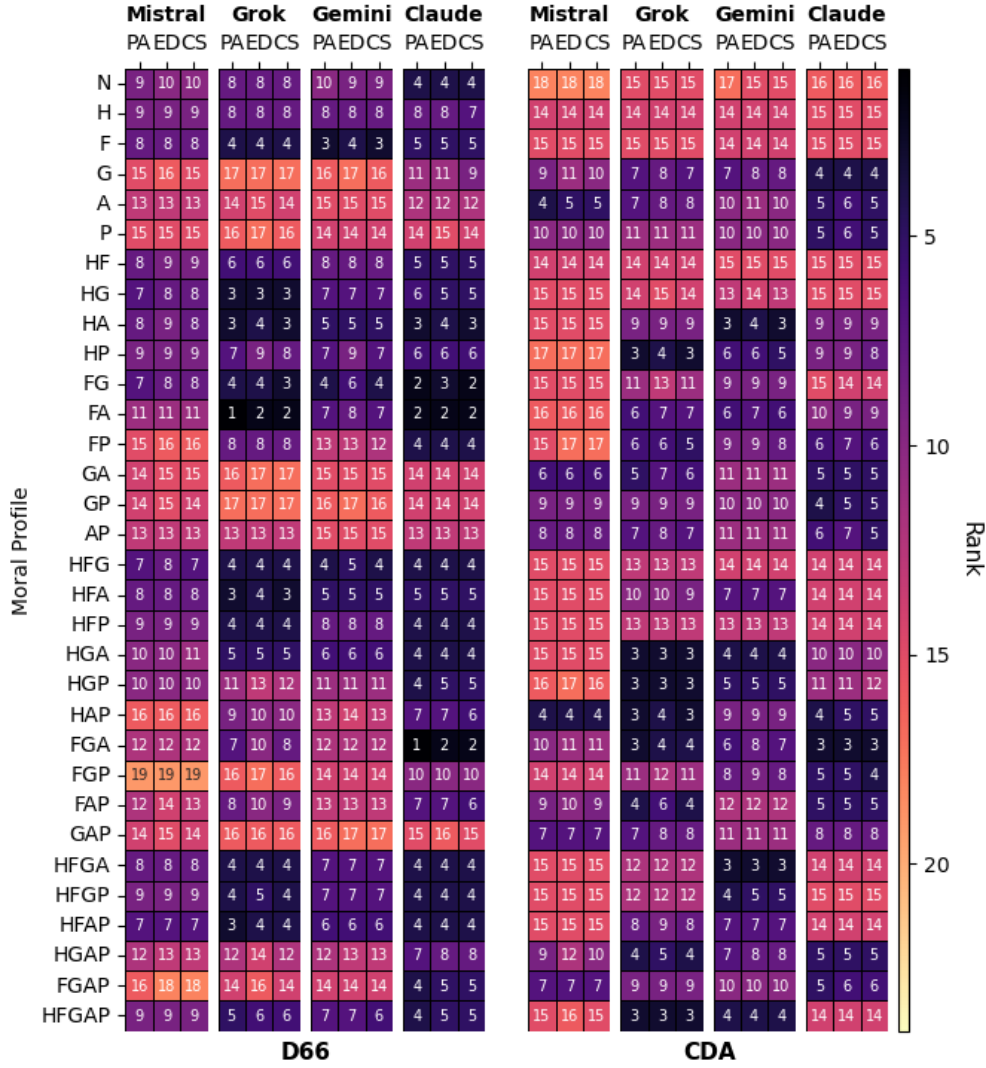


Figure 6: Assigned ranks for D66 (left) and CDA (right) calculated for each metric for all profiles in all models. The colouring is done on the resulting rank.

Both D66 and CDA are political parties holding considerable amount of seats in the Dutch parliament, yet their representation in the top three rankings of the profiles by the LLMs is limited. Hence, extra inspection for these two parties is done in the heatmaps of Figure 6, where for all models, profiles, and metric the rank for both parties is given as the colour scale. The rankings for D66 in the four leftmost columns shows the party falls just outside of the top three very often, especially in Claude where they are often ranked fourth. Mistral deviates from this observation as D66 is never seen within the top six. Great differences in rank for certain profiles between models are seen, such as for the profiles FA and FP. Those differences in rank between models for the same profile are also observed for CDA in the four rightmost columns. Again Mistral is the model behaving differently; CDA often has a higher rank in this model than in the others for the same moral priming. For both parties, differences in assigned rank due to the metric is minimal.

5.5 Diffusion of Voting Behaviour

The response vectors of all participating parties and all 640 simulated voters are visualised into a 2-dimensional space by the t-SNE algorithm. In Figure 7 the t-SNE plot is presented where the datapoints are grouped in four 'clouds' each representing one model and their distribution of generated datapoints. The political parties seem to group into a few clusters. The models overlap with most clusters as their spread over the space is broad. One cluster at the top-right corner of the plot is not captured by any of the model clouds; the voting behaviour of the models do not appear to overlap with the political parties in this cluster.

Mistral shares very little overlap with other models and the clusters of the Dutch parties, though it is the only model overlapping with DENK. Claude's behaviour results in three clusters of points spread over the space. Gemini and Grok overlap greatly with each other, though differences in the density of their distributions are seen. All models have a widespread distribution over the space, often overlapping with each other and the clusters formed by the parties.

Figure 8 presents the same t-SNE plot as previously but now the each clouds represent the distribution of voters that were primed by only one of the five moral foundations and the voters that were not primed. They appear to divide into three clusters: the neutral personas, individualising foundations, and binding foundations. The neutral personas gather around the cluster of political parties at the top-right. Harm/Care and Fairness/Reciprocity crowd together at the mid-right side of the plot, overlapping a bit with each other. Together they overlap mostly with the cluster of political parties on this side of the figure. In-group/Loyalty, Authority/Respect, and Purity/Sanctity overlap mostly on the top left-side of the plot with another cluster of political parties. Only for In-group/Loyalty is part of its distribution of datapoints found not overlapping with the other foundations. This distribution is wider and included more of the parties. In the middle of the plot the parties NSC, ChristenUnie, and CDA can be seen sharing no overlap with any of these profiles.

The models appear to behave differently between the two sets of foundations as well as the neutral persona. Voting behaviour of the models for each of these groups appear to share similarities to different sets of political parties. Section 5.1 presented differences in the share of consistency between foundations. However, none of the foundations were seen randomly spread over the space. This implies the models display similar voting behaviour for the same morally primed profile and therefore associates the underlying moral values with a set of opinion. Though, the size of this set differs across profiles.

The full distribution of for presence of each foundation in a profile are illustrated in the sub-figures of Figure 9. The models voting behaviour appears to spread for addition of more morals; the datapoints do not position solely in the area of a single foundation for multiple morals.

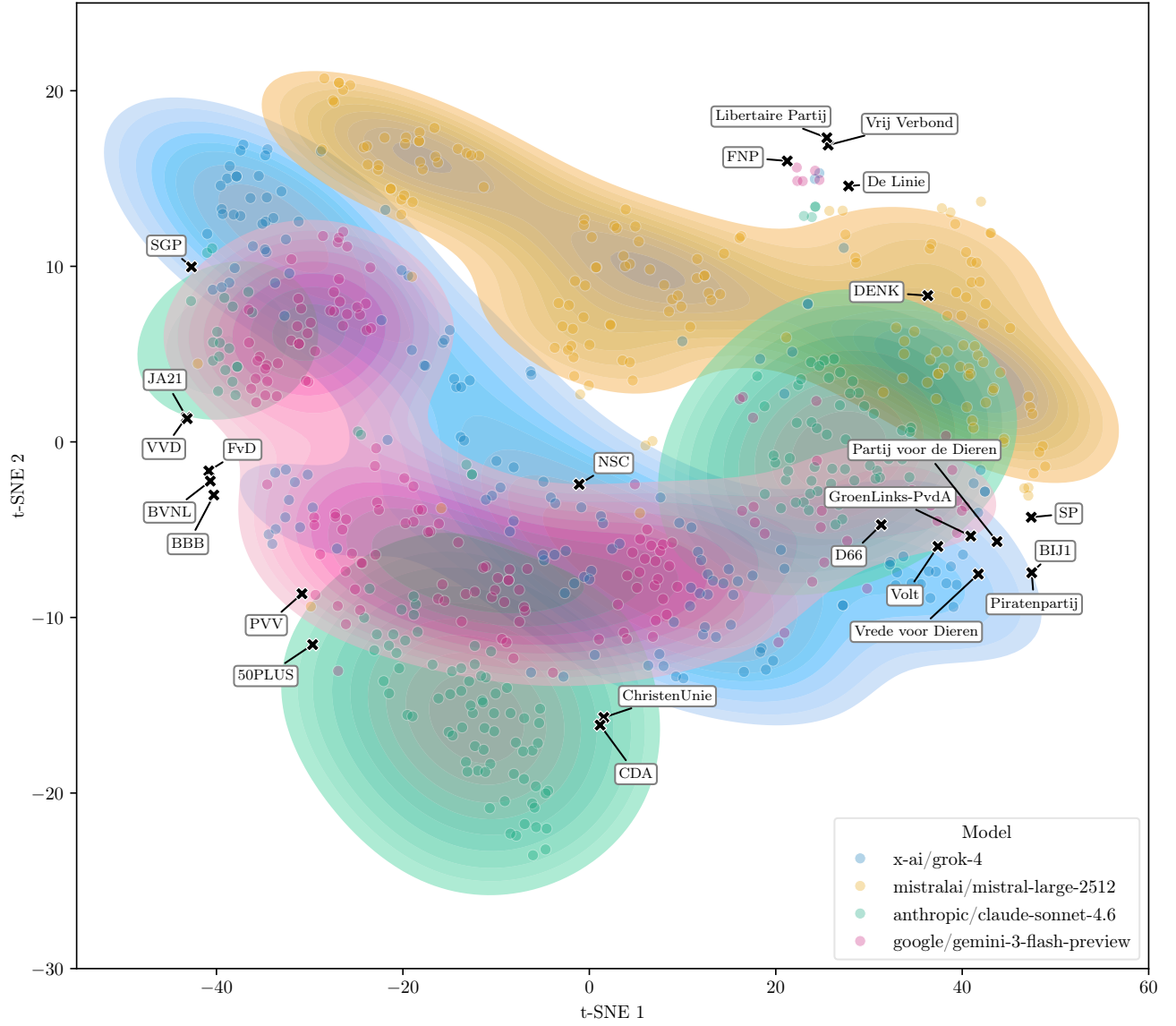


Figure 7: t-SNE projection of all political parties and model behaviour for all voter profiles. Coloured regions represent the kernel density estimate of voter distribution for each of the four models, with the density threshold set at 0.3

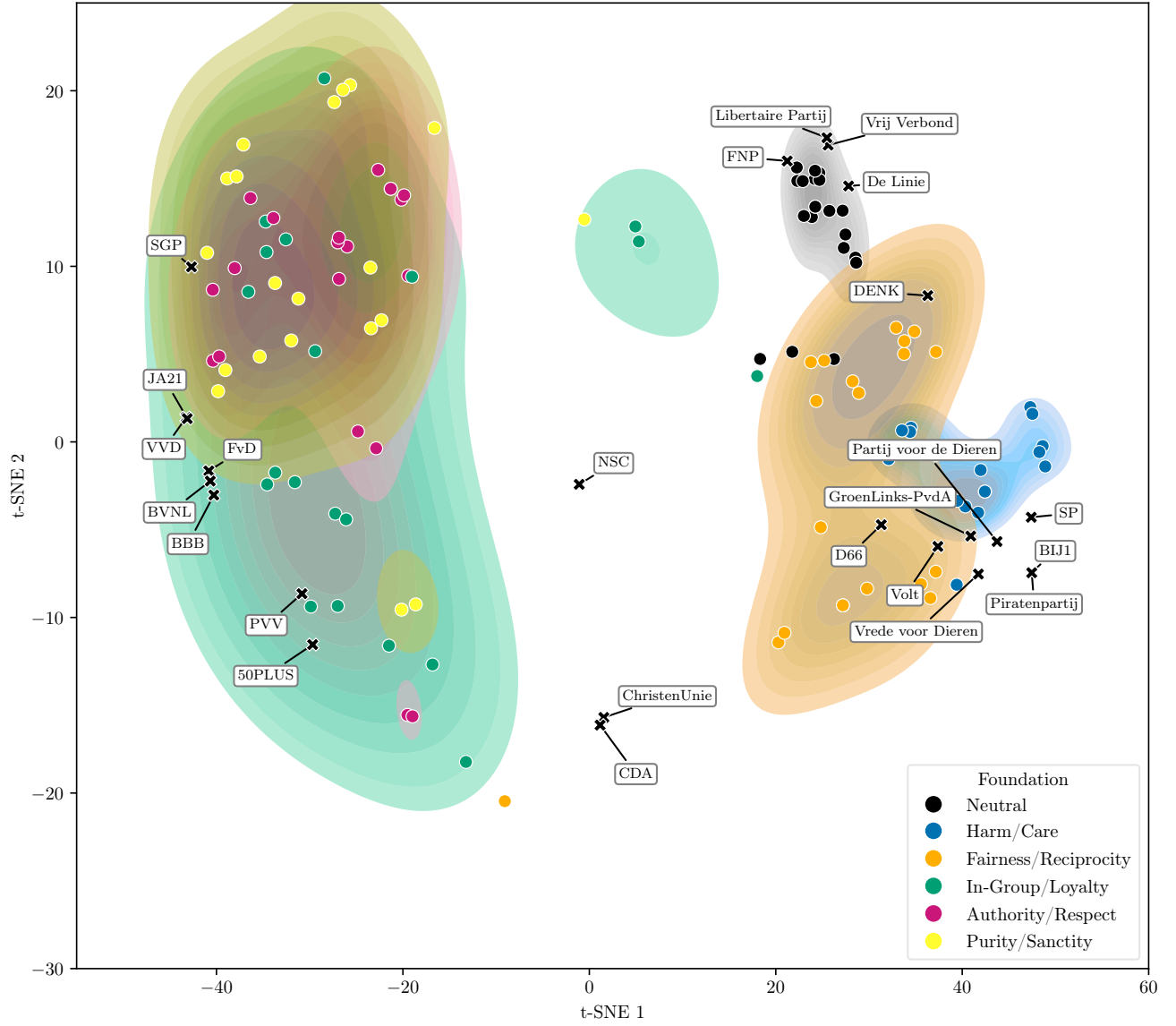
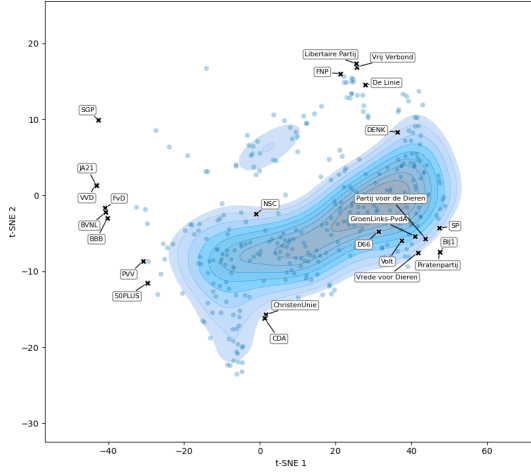
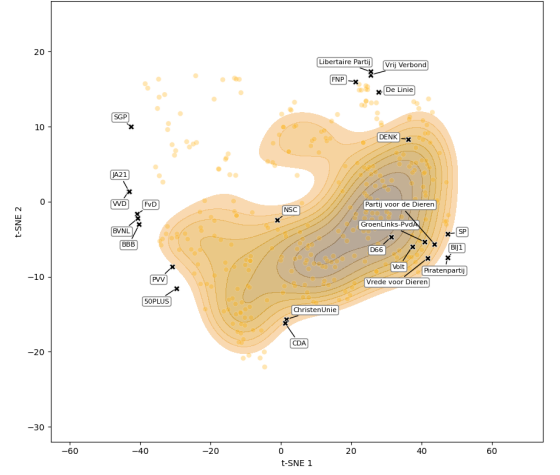


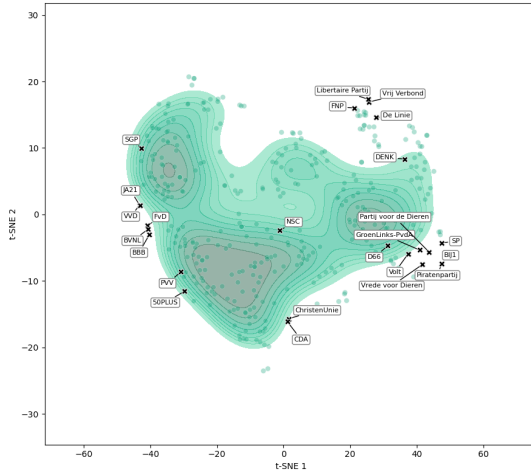
Figure 8: t-SNE projection of all political parties and model behaviour for all voter profiles containing exclusively a single moral foundation. Highlighted regions represent the kernel density estimate of voter behaviour, with the density threshold set at 0.3, across models.



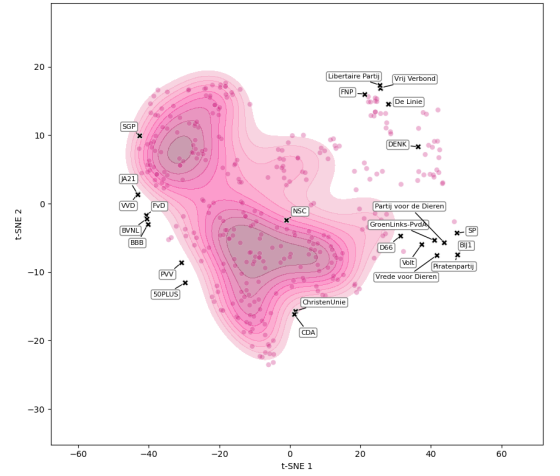
(a) Harm/Care



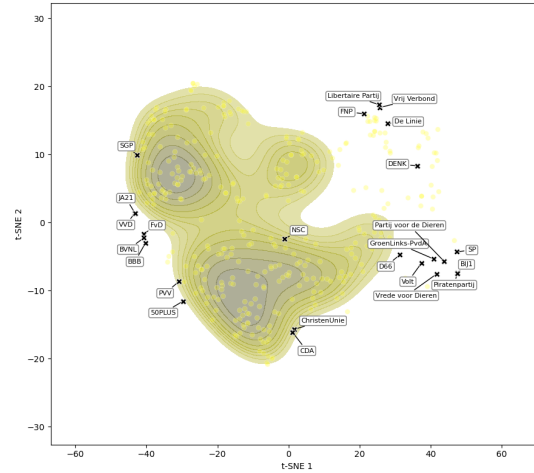
(b) Fairness/Reciprocity



(c) In-group/Loyalty



(d) Authority/Respect



(e) Purity/Sanctity

Figure 9: t-SNE projections of all voter profiles and political parties, with each sub-figure isolating presence of one moral foundation: (a) Harm/Care, (b) Fairness/Reciprocity, (c) In-group/Loyalty, (d) Authority/Respect, (e) Purity/Sanctity. Shaded regions show the kernel density estimate of voter behaviour across models, with the threshold set at 0.3.

6 Discussion

6.1 Interpretation of Results

The results of the experiments suggest that voting behaviour of LLMs on political statements was influenced when prompted on behalf of a morally framed voter. None of the models focused on a specific political party nor a cluster of parties in the projection of the t-SNE algorithm. Indeed, a broader scattering of datapoints was observed for the primed voters than for the political parties. Previous findings claimed that models behaviour resulted in a polarising view of the Dutch political landscape, our findings do not support this as the models appear to be steerable. Those previous results might stem from the voting behaviour of the political parties themselves as many of them often vote similar, resulting in narrow distribution. It cannot be expected from LLMs to then perfectly distinguish the parties purely from their responses to the statements of StemWijzer.

The models took a position on political issues for voters more often than not. Claude answered neutral most frequently and Grok the least which may be a contributing factor in their respective consistency scores. Previous literature found the alterations of prompts to be somewhat successful for gaining a response on political statements from models who previously refused to do so. We have not done so as obtaining a stance was not the purpose of this research.

The metrics differed somewhat in their rankings of the political parties, where the method of StemWijzer, Percentage Agreement, was favourable for some parties. A factor of influence is the punishment of a neutral answer when it is not in agreement with the answer of the political party as this metric does not preserve the ordinal scale of the possible responses. This is taken into account by Euclidean Distance and Cosine Similarity leading less ties in the top three rankings. However, by these metrics the neutral personas sometimes have a stronger match to their rankings than morally primed personas although not that much political orientation would be given from their answers. Euclidean Distance and Cosine Similarity might be good alternative metrics to Percentage Agreement when a neutral answer is given.

One cluster of political parties was exclusively covered by the models for the neutral personas. An explanation could be that those parties have a less coherent political identity as only the voters that do not express opinion share their observed position. They might be seen as close to the algorithm due to lack of information rather than shared features. This is supported by the top three rankings of the neutral persona, where only one of the parties in this cluster is seen.

GroenLinks-PvdA was most frequently observed in the top three highest aligned parties for all metrics. This party was near always closest to a voter if it was primed with exclusively one of or inclusion of both the individualizing foundations Harm/Care and Fairness/Reciprocity. An exception was observed for Mistral which had the least counts for GroenLinks-PvdA. Instead, Mistral's most occurring party was PVV which was less commonly seen on the high-end of the rankings in other models. Interestingly, these two parties were observed in different clusters of parties. Mistral shared the least overlap with other models, appearing to share less similarities in their voting behaviour.

The data showed some variations in the consistency of the profiles primed with exclusively one the moral foundations. Because the foundations themselves are freely interpretable, some variance is expected. The differences may stem from the way these models embed, and consequently interpret the morals. The less consistent the profile, the bigger the set of opinions the model associates the corresponding morals with. Still, the findings suggest a set space for all foundations where the

binding foundations group together as well as the individualizing foundations. Both in the area of different clusters of political parties suggesting an interpretation of these morals that corresponds to the voting behaviour of the aligned parties.

The underlying reasons that result in the voting behaviour of each of the LLMs cannot be explained from our work. A possible explanation could be refusal to be either negative or positive on certain policy issues, as found in previous literature. The extent of fine-tuning of the models such that they do not express opinion on sensitive issues is unclear.

6.2 Limitations

Several limitations of this study should be noted. Due to the method of generating the morally primed profiles, the order in which the foundations were implicitly mentioned was consistent (e.g. if Purity/Sanctity was present in the persona it would always be seen last). LLMs have been shown to be sensitive to the input order of their prompt [10], this might have made the models weigh the foundations that were mentioned first more heavier in their decision. For that reason, we recommend counterbalanced sub-prompt ordering for future research.

The theorem upon which the profiles are constructed provides a base for moral values. Political decision making can be more nuanced and might involve more values, especially in a multi-party democracy. How citizens actually interact with LLMs when asking for voting advice or when consulting for help when filling in a VAA is unknown. We propose our method of a voter explaining their opinion on the basis of moral values to be a natural way of interaction with a model. However, there is no previous research supporting this idea. We did not permit further explanation or interaction, which limits behaviour of the LLMs.

A limitation regarding assessment of political alignment is found in the three-point answer scale of StemWijzer. This leaves limited space for nuances regarding opinion on the statement and the political parties very rarely answer neutral on a statement. We did not increase the scale to, for example, five-points due to the then arising difficulties when comparing the models voting behaviour to the voting behaviour of the political parties. The VAA Kieskompas could provide a solution to this issue as it uses a five-point scale. However we did not use this VAA because of the undisclosed determined position in the political dimensions of each statement, which is used to map voters onto the political landscape.

7 Conclusion and Future Research

This thesis examined the effect of moral framing in voter prompts on the responses of LLMs when asked to fill in the VAA StemWijzer for the Dutch general election of 2025, assessed by comparison against the positions of participating political parties on the statements. Morally primed voter profiles were generated on the basis of the Moral Foundation Theory in a first-person perspective, after which four models were asked to answer on all statements of the VAA on each voters behalf. All interaction with the models was done in Dutch and none of the parties names were ever disclosed to the models.

Q1. Are LLMs consistent in their response to political statements on behalf of voters?

Generally, the LLMs did respond with a stance more often than not, though differences in neutral position were found which could explain the differences in their consistency. Furthermore, the consistency of models response varied between the moral foundations and combinations thereof present in the voter profile.

Q2. Does the moral complexity of a voter profile affect the consistency in answers on the VAA of a model?

No significant relationship between the moral complexity of a voter profile and the consistency was found.

Q3. To what extent does choice of metric to measure alignment affect the ranking of political parties to LLMs response?

Choice of metric affected the frequency with which certain political parties appeared amongst the best-ranked parties; though the rankings for the each profile averaged over versions resulted in largely the same outcomes. A difference in the strength of match to the parties between voter profiles was observed for different metrics.

Q4. How do individual moral foundation or combinations of moral foundations influence the strength of match between LLM and political parties response?

Stronger matches to highest-ranked political parties were observed for the exclusive and co-occurring presence of the moral foundations Harm/Care and Fairness/Reciprocity in a voter persona.

Q5. Do models capture the moral value structures underlying Dutch political ideologies when answering on behalf of the morally primed voter profiles?

By examination of a t-SNE projection, the models voting behaviour resulted in a broad dispersion of datapoints covering multiple clusters of political parties.

Models voting behaviour on political issues is influenced when responding on behalf of a morally primed voter. Some morals or combinations thereof led to more consistent response and stronger matches to parties, and some parties were observed ranked highly more frequently. However, many parties were observed close together in their own voting behaviour, forming a less broad dispersion than the models. Thus, models response to the moral frames resulted in a richer opinion landscape than the political parties.

To address some of the limitations identified, future research could explore alternate methodological choices. Firstly, an assessment method that allows for greater nuances, such as the VAA Kieskompas. If, upon inquiry, they are willing to disclose their determined positions of each statement on the political spectrum, research of models behaviour towards policy specific issues could be considered. Secondly, limited research has been done on how voters would consult a LLM on voting advice. Research on this interaction might present alternative value frames on which the models behaviour can subsequently be investigated. Lastly, the observed compact spread of

the twenty-four political parties in their response behaviour could be further examined within political science. It could be that the statements of StemWijzer do not allow for full separation of the parties.

References

- [1] Salah Feras Alali, Mohammad Nashat Maasfeh, Mucahid Kutlu, and Saban Kardas. Measuring political stance and consistency in large language models. In Ricardo Campos, Adam Jatowt, Yanyan Lan, Mohammad Aliannejadi, Christine Bauer, Sean MacAvaney, Avishek Anand, Zhaochun Ren, Suzan Verberne, Nan Bai, and Masoud Mansoury, editors, *Advances in Information Retrieval*, pages 442–457, Cham, 2026. Springer Nature Switzerland.
- [2] Rudy B. Andeweg, Galen A. Irwin, and Tom Louwerse. *Governance and Politics of the Netherlands*. Red Globe Press, 5th edition, 2020.
- [3] Directie Coördinatie Algoritmes (DCA) Autoriteit Persoonsgegevens. Vervolgonderzoek AI-chatbots als stembulp: AI-chatbots als stembulp bij gemeenteraadsverkiezingen. <https://www.autoriteitpersoonsgegevens.nl/actueel/ai-chatbots-negeren-lokale-partijen-bij-stemadvies>, March 2026. Accessed: 2026/04/19.
- [4] Tanise Ceron, Neele Falk, Ana Barić, Dmitry Nikolaev, and Sebastian Padó. Beyond prompt brittleness: Evaluating the reliability and consistency of political worldviews in llms. 2024.
- [5] Ernesto de León, Fabio Votta, Theo Araujo, and Claes de Vreese. 1 in 10 Dutch citizens are likely to ask AI for election advice. This is why they shouldn’t. <https://www.uva.nl/en/shared-content/faculteiten/en/faculteit-der-maatschappij-en-gedragwetenschappen/news/2025/10/dont-ask-ai-for-election-advice.html>, October 2025. Accessed: 2026-04-19.
- [6] Shangbin Feng, Chan Young Park, Yuhao Liu, and Yulia Tsvetkov. From pretraining data to language models to downstream tasks: Tracking the trails of political biases leading to unfair nlp models, 2023.
- [7] Jillian Fisher, Shangbin Feng, Robert Aron, Thomas Richardson, Yejin Choi, Daniel W Fisher, Jennifer Pan, Yulia Tsvetkov, and Katharina Reinecke. Biased ai can influence political decision-making. 2024.
- [8] Thomas Fossen and Joel Anderson. What’s the point of voting advice applications? competing perspectives on democracy and citizenship. *Electoral Studies*, 36:244–251, 2014.
- [9] Jesse Graham, Jonathan Haidt, and Brian A Nosek. Liberals and conservatives rely on different sets of moral foundations. *Journal of personality and social psychology*, 96(5):1029, 2009.
- [10] Bryan Guan, Tanya Roosta, Peyman Passban, and Mehdi Rezagholizadeh. The order effect: Investigating prompt sensitivity to input order in llms. 2025.

- [11] Jonathan Haidt and Craig Joseph. Intuitive ethics: How innately prepared intuitions generate culturally variable virtues. *Daedalus (Cambridge, Mass.)*, 133(4):55–66, 2004.
- [12] Jonathan Haidt and Craig Joseph. The moral mind: How five sets of innate intuitions guide the development of many culture-specific virtues, and perhaps even modules. *The Innate Mind, Volume 3, Foundations and the Future*, 3, 04 2008.
- [13] Jochen Hartmann, Jasper Schwenzow, and Maximilian Witte. The political ideology of conversational ai: Converging evidence on chatgpt’s pro-environmental, left-libertarian orientation, 2023.
- [14] Naomi Kamoen and Bregje Holleman. I don’t get it: Response difficulties in answering political attitude questions in voting advice applications. In *Survey Research Methods*, volume 11, pages 125–140. European Survey Research Association, 2017.
- [15] Autoriteit Persoonsgegevens. Ap waarschuwt: chatbots geven vertekend stemadvies. <https://www.autoriteitpersoonsgegevens.nl/actueel/ap-waarschuwt-chatbots-geven-vertekend-stemadvies>, October 2025. Directie Coördinatie Algoritmes (DCA).
- [16] Yujin Potter, Shiyang Lai, Junsol Kim, James Evans, and Dawn Song. Hidden persuaders: Llms’ political leaning and their influence on voters. 2024.
- [17] Luca Rettenberger, Markus Reischl, and Mark Schutera. Assessing political bias in large language models. *Journal of Computational Social Science*, 8(2), Feb 2025.
- [18] Gabriel Simmons. Moral mimicry: Large language models produce moral rationalizations tailored to political identity. 2022.
- [19] Laurens van der Maaten and Geoffrey Hinton. Visualizing data using t-sne. *Journal of Machine Learning Research*, 9(86):2579–2605, 2008.

Appendix A Prompts for Profile Generation

```
1 class MoralProfiles(dspy.Signature):
2     """
3     Mensen hebben volgens de Moral Foundation Theory (MFT) vijf aangeboren
4     morele intuïties die verschillen per cultuur en per individu. Dit zou
5     verschillen in morele redenering verklaren en daardoor ook van invloed zijn
6     in mensen hun politieke houding. De theorie beschrijft de volgende vijf
7     morele intuïties:
8
9     Schade/Zorg: De gevoeligheid voor lijden en pijn die ertoe leiden dat
10    mensen willen zorgen voor de kwetsbare individuen van de samenleving door
11    ze te beschermen tegen gevaar.
12    Eerlijkheid/Wederkerigheid: De wil voor rechtvaardigheid, eerlijkheid en
13    wederkerigheid zodat valsspelen niet wordt getolereerd.
14    Groep/Loyaliteit: Loyaliteit en prioriteit aan de eigen groep waar
15    saamhorigheid en zelfopoffering zwaar wegen.
16    Autoriteit/Respect: Sociale orde en stabiliteit door respect voor
17    rechtvaardige leiders en de volbrenging van de aangewezen rol in
18    de maatschappij
19    Zuiverheid/Heiligheid: Het bewaken van reinheid en ontwijken van
20    vervuiling van het lichaam maar ook van de sociale orde. Dit is op zowel
21    een fysieke als een symbolische manier.
22
23    De intuïties Schade/Zorg en Eerlijkheid/Wederkerigheid worden
24    geclassificeerd als “individuele fundaties” die het welzijn en recht
25    van een individu hoog houden. De intuïties Groep/Loyaliteit,
26    Autoriteit/Respect, en Zuiverheid/Heiligheid worden benoemd als
27    “bindende fundaties” door hun promotie van maatschappelijke cohesie,
28    plicht en zelfcontrole.
29
30    Schrijf een korte beschrijving van een persoon die niet weet wat die moet
31    stemmen maar wel weet welke moralen belangrijk zijn voor hun politieke
32    houding:
33    - Laat het natuurlijk klinken
34    - In de beschrijving moet impliciet duidelijk zijn welke moralen belangrijk
35    zijn door uitingen van zorgen, prioriteiten of voorbeelden.
36    - Schrijf 5 zinnen of minder verteld in de ik-vorm
37    - Noem hierbij de MFT of de intuïties niet
38    - Gebruik verschillende voorbeelden, situaties of formuleringen
39    """
40    morals = dspy.InputField(desc="Gebruik deze moralen voor de beschrijving,
41    laat ze even zwaar wegen en noem de andere niet")
42    variation = dspy.InputField()
43    voter_profile = dspy.OutputField(desc="Profiel van een Nederlandse
44    stemgerechtigde in de ik-vorm")
45
```

Listing 1: The `dspy.Signature` for generation of the morally primed voter profiles.

```

1 class NeutralVoter(dspy.Signature):
2     """
3     Schrijf een korte, neutrale beschrijving van een persoon die niet weet
4     wat die moet stemmen. De desbetreffende persoon heeft geen duidelijke
5     politieke houding.
6
7     - Laat het natuurlijk klinken
8     - Schrijf maximaal 2 zinnen verteld in de ik-vorm
9     """
10    variation = dspy.InputField()
11    voter_profile = dspy.OutputField(desc="Profiel van een Nederlandse
12    stemgerechtigde in de ik-vorm")
13

```

Listing 2: The `dspy.Signature` for generation of neutral voter profiles.

Appendix B Prompt for Filling in the VAA

```

1 class VoterQuestion(dspy.Signature):
2     """
3     Je helpt mensen met politieke stellingen te beantwoorden namens hun.
4     Meerdere stellingen worden aan je gepresenteerd, beantwoord deze namens het
5     profiel van de kiezer.
6     Antwoord elke stelling met eens, neutraal, of oneens.
7     Geef precies één antwoord per vraag, begin elk antwoord op een nieuwe regel
8     en nummer de antwoorden niet.
9     Geef geen verdere verklaring, tekst of herhaling van de stelling.
10    """
11    voter_profile: str = dspy.InputField(desc="Beantwoord namens deze
12    kiezer")
13    statements: str = dspy.InputField(desc="Stellingen")
14    answers: str = dspy.OutputField(desc="Beantwoord elke stelling met eens,
15    neutraal, of oneens")
16

```

Listing 3: The `dspy.Signature` to let the model fill in the voting advice application.

Appendix C Abbreviation Political Parties

Name	Abbreviation
GroenLinks-Pvda	GL-PvdA
ChristenUnie	CU
Partij voor de Dieren	PvdD
Piratenpartij	PPNL
Vrij Verbond	VV
Vrede voor Dieren	VvD
Libertaire Partij	LP
50PLUS	50+
De Linie	DL

Table 5: For illustration purposes some parties with longer names have been abbreviated in tables. The left column states the name of the party as referenced to in the StemWijzer for the general election of 2025 and the right column their assigned abbreviation. Parties not seen in this table have not been altered in their name.