

(Note: This is the first draft of an investigation that is very much in progress. We prioritized soundly discussing the original method we introduced; the sections about its application for empirical analysis are a bit more schematic and preliminary. We look forward to hearing your comments, feedback and suggestions about how to move this project forward!)

APPLYING BIPOLAR CLASS ANALYSIS TO ANALYZE THE DEVELOPMENT OF POLITICAL OPPOSITION VECTORS IN THE US

(DRAFT PREPARED FOR MESS 2024; PLEASE DO NOT CIRCULATE WITHOUT PERMISSION)

MANUEL CUERNO, FERNANDO GALAZ-GARCÍA, SERGIO GALAZ-GARCÍA,
AND TELMO PÉREZ-IZQUIERDO

ABSTRACT. This investigation empirically studies the number, opinion constraint structure, and relative size of "vectors of political opposition"—the best known of which is the liberal-conservative divide—in the American electorate from the 1980s up until the present day. For this task we develop and use Bipolar Class Analysis (BCA), a method that uses survey data to estimate vectors of political opposition (VPOs) by clustering respondents according to how similar their answers to opinion questions toggle back and forth between support and rejection values. We discuss how BCA tackles inferential issues observed for similar existing methods (like Relational Class Analysis, or RCA) and show that it produces more accurate estimations of VPOs. We apply BCA to estimate VPOs using ANES data from 1988 to 2020. Results show that liberal-conservative divide has become the dominant vector of position in the American electorate, and that this process has been organized in two stages: one in which this vector lost population but increased constraint across issues opinions along a left-right wing divide, and another, observed between 2016 and 2020, when opinions became less constrained but encompassed a larger number of people.

. (M. Cuerno) DEPARTMENT OF QUANTITATIVE METHODS, CUNEF UNIVERSIDAD, SPAIN
(F. Galaz-García) DEPARTMENT OF MATHEMATICAL SCIENCES, DURHAM UNIVERSITY, UNITED KINGDOM
(S. Galaz-García) SOCIAL SCIENCES DEPARTMENT / INSTITUTO JUAN MARCH, UNIVERSIDAD CARLOS III-MADRID, SPAIN
(T. Pérez-Izquierdo) DEPARTMENT OF ECONOMIC ANALYSIS, UNIVERSITY OF THE BASQUE COUNTRY, SPAIN
. *E-mail addresses:* manuel.mellado@cunef.edu, fernando.galaz-garcia@durham.ac.uk, sergio.galaz@uc3m.es, telmo.perezizquierdo@ehu.eus.
. *Date:* October 14, 2024.
. *Key words and phrases.* Quantitative Methodology, construal estimation, Relational Class Analysis, Correlational Class Analysis, Bipolar Class Analysis, Simulation Analysis .

SECTION 1. INTRODUCTION

This investigation empirically analyzes the number, structure of opinion constraint, and relative size of vectors of political opposition in the American mass electorate since the 1980s up until the present day. By "vectors of political opposition", or VPOs, we refer to arrays of political stances that, by sharing complementary interpretations of what is desirable and undesirable across political issues, have a common frame of reference to interpret what is at stake in political process. The best known VPO is the left-wing of "lib-con" divide, which draws a line between liberals with left-wing positions across all political issues and conservatives with right wing positions across them.

We seek to contribute to filling a gap in the the burgeoning literature in the subfield of public opinion studying how people organize their political beliefs. These studies have focused in the studying political polarization or how strong people's opinions organize around the well known standard liberal-conservative divide. But we still have few general examinations of the *structure* of political opposition vectors that describe political opinions in the US. We count with little information on whether there other vectors of political opposition other than the standard "lib-con" divide, how recurrent these forms might be among US citizens, and the specific patterns of attitudinal heterogeneity around which they organize political positions. Findings on these matter carry important implications for electoral outcomes and democratic governance against the backdrop of the increasing lib-con polarization among political parties and elites in the US (Robert Lupton and Thornton 2015; Shor and McCarty 2018; Poole and Rosenthal 2007). In this context, the presence of vectors of political opposition other than the left-right wing divide is likely to induce electoral volatility and suboptimal legislative and public policy outcomes.

Public opinion findings currently offer three stylized expectations on how VPO structures have historically developed in the US. One of them contends that the lib-con vector has become stronger and more dominant over time (Hare 2022; Abramowitz and Saunders 2014). The second take expects that the number, correlational characteristics and size of VPOs should have remained stable over time. Another view holds that, at the mass level, the political opinions of US citizens have maintained stances on issues that have very few relationships with one another (Morris Fiorina and Pope 2006; Converse 1964). This view then posits that the American electorate has consistently ranked as a member a zero-degree, or "null" vector of political opposition. Another take argues that VPOs in the use have been characterized by a "tripartite stability". Drawing from the findings of Baldassarri and Goldberg (Baldassarri and Goldberg 2014), this view contends that the organization of political opinions among American citizens has remained divided into three VPOs of roughly

equal size: the standard lib-con divide, a "null" vector of political opposition, and an "alternative" vector that draws a line between people who are culturally liberal but economically conservative with others who are culturally traditional but economically progressive .

Recent works have started to produce cumulative evidence supporting the lib-con strengthening hypothesis. However, their findings are not conclusive because the item-response methodological framework they apply in their investigation is commonly catered to assess how well people's opinions organized a unique, latent political opposition vector (a latent vector that is assumed, but often not explicitly shown align to a pure left-right wing divide) (fowler'etal'2022 ; Hare 2022). However, they cannot explore whether there exists groups of people that bundle their political opinions in ways that could be better characterized as adhering to alternative political belief systems (Converse 1964). To our knowledge, the only study on the organization of political opinions at a mass level directly catered at studying diversity of political opposition vectors is Baldassarri and Goldberg's (2014). Their findings, which form the empirical basis for the "tripartite stability" described in the preceding paragraph, are based in the use of Relational Class Analysis (or RCA), an innovative respondent partitioning method that has been recently developed within sociological research (Goldberg 2011). In a nutshell, this method analyzes survey answers to attitudinal questions to generate measures of response pattern similarity across pairs of survey takers, assembles a weighted response similarity graph using these measures (Newman 2006) and applied community detection algorithms to divide this similarity graphs into clusters of respondents sharing in common specific forms of response pattern similarity.

RCA represents an important methodological achievement that has allowed to generate mappings of recurrent patterns of political opinion heterogeneity without assuming that this heterogeneity is best described as organized along a unique line of political opposition. In this investigation we continue developing the line of methodological development that RCA opened by noting several important opportunities for amelioration. We find that RCA's ability to generate reliable and accurate estimations of political opposition vectors is compromised by its analysis of respondents' answers as bounded cardinal variables, on one hand, and on the other—and most importantly—by how it computes response pattern similarity among respondents. We discuss how these two characteristics lead to inferential problems for the analysis of typical data structures in public opinion surveys.

To tackle these problems, we introduce Bipolar Class Analysis, or BCA, as a next-generation respondent partitioning method. BCA uses a "polar" response similarity among survey takers that captures how similarly their answers toggle back and forth between opinions expressing support or

rejection to question statements. In the main body of this paper draft we present an intuitive exposition of BCA’s main tenets; we also offer a more formal introduction of the method in Appendix A. We also conduct simulation analyses to evaluate the performance of RCA as an accurate method for the estimation of PVOs relative to BCA using simulation analyses. Their results strongly indicate that BCA is better capable of accurately estimating the number, relative size, and specific structure of constraint across opinions than RCA.

We apply BCA to analyze American National Election Survey (ANES) data in order to assess how VPOs have historically developed in the us between the 1980s up until 2020. We recognize temporal comparability challenges presented by the high irregularity with which public opinion items have been included across ANES waves. We address this issue by identifying sets of public opinion questions that were consistently included in the ANES waves with the highest density of this type of items in the 1980s, 1990s 2000s, and the three waves that were conducted in the 2010s (2012, 2016, 2020). This strategy allows tackling compositional heterogeneity issues while also allowing to explore how robust findings are to issue selection.

Broadly speaking, we find evidence supporting that lib-con political opposition vectors have become stronger over time. Our results suggest that, in the late 1980s, the way Americans bundled together their political opinions did not conform to a lib-con configuration. This results is in stark contrast to the one we find for 2020. In this year, every survey respondent ranked as member of a VPO that exhibited lib-con constraint structures. We also find that this change did not evolve in a linear nor qualitatively uniform fashion. Our results suggest it can be decomposed in two stages: in the first, spanning between the 1980s and the 2016 lib-con forms of political opposition underwent a process of *concentration*: while constraints between opinions aligned to a standard left-right wing understanding of politics became stronger, the number of people that adhered to these opinion configurations shrank; in the second, which took place much more recently (2016-2020), lib-con VPOs exhibited a process of *generalization*: while left-right ideological constrain shrank but their relative size increased.

[OK] The structure of this draft proceeds as follows. In the first section, we discuss and critically assess the use of RCA to estimate VPos. In the next section we introduce Bipolar Class Analysis and evaluate its accuracy as a VPO estimation method relative to RCA. The third section describes the research design and analytical strategy to evaluate the historical development of VPOs in the US. The subsequent section reports the results of our analysis.

	<i>Procedure</i>	<i>Description</i>
1.	Numerization of answers	Respondents' answers are recoded into cardinal numeric variables with a range of [0,1]; zero- and one-values represent farthest left and right-wing positions.
2.	Computation of baseline measures of response similarity among pairs of survey-takers	Measure of similarity: <i>relationality</i> Mean of scores of answer similarity between pairs of questions with range of [0,1] Pair-wise answer scores measure similarity in terms of distance and direction in of answer changes.
3.	Assemblage of similarity graph	A weighted graph of answer similarity across respondent is built. Respondents are represented as nodes and similarity scores among them as ties.
4.	Noise reduction and clean-up of similarity graph	Non-significant ties are dropped from the graph; the remaining ones are transformed into absolute values and centered by mean relationality.
5.	Division of graph into clusters	Graph is split into clusters of respondents with common patterns of response similarity using spectral community detection algorithms (Newman 2006; Blondel 2011).

FIGURE 1. Procedural algorithm for Relational Class Analysis (RCA)

Notes: The description of procedures is drawn from Goldberg's methodological introduction of RCA (2011, 1403-1410 and Appendixes A and E). Subsequent applications in public opinion research have used modified versions of RCA (Baldassarri and Goldberg 2014; Brensinger and Sotoudeh 2022); we describe this method in its original version since it remains its most extensively documented and readily available iteration.

SECTION 2. RCA AS A METHOD FOR THE DETECTION OF POLITICAL OPPOSITION VECTORS: STRENGTHS AND LIMITATIONS

RCA (and later methods for construal estimation subsequently developed) estimates vectors of political opposition from survey data using an algorithm of procedures set by Goldberg (Goldberg 2011). This algorithm consists of 5 steps: answer numerization; computing response similarity between pairs of respondents, assemblage of a similarity graph across respondents, applying clean-up and noise reduction procedures to this graph, and finally, dividing into clusters of respondents with common response patterns through the use of community detection algorithms. These clusters are then taken as an estimation of the number, specific patterns of association between opinions, and relative population of VPOs at a societal level.

In a seminal study, Baldassarri and Goldberg (2014) analyzed ANES survey data to study how the number, associations between opinions, and size of VPOs developed in the US from 1984 to 2004. They found that, at the mass level, Americans rank as members of one of three equally populated VPOs of roughly equal size. The first vector of political opposition corresponds to the "standard left-right" wing divide that puts people who are culturally and economically liberal in opposition with fiscal and cultural conservatives. Thus, when answers values would be reordered such that smaller values corresponded to increasingly left oppositions, and larger values to increasingly right ones, this line of opposition would be characterized by positive associations across opinions in a way similar to the correlation structure presented in the left chart of Figure 2. A second VPO, referred to by Baldassari and Goldberg with the simple rubric of "alternative" , draws a line of opposition between people who could be referred to as "fiscally conservative progressives" who are economically conservative but culturally liberal and "fiscally redistributive conservatives". As shown by panel 2, the correlation structure across opinions in this VPO is thus negative between economic and cultural issues. The third and final VPO that Baldassarri and Goldberg identify is a "null" VPO that groups together people described by Converse (1964) and Fiorina (Morris Fiorina and Pope 2006) who hold issue positions largely dissociated from one another.

The introduction of RCA marked a methodological breakthrough in our capacity to generate an assumption-free, ground up mapping of vectors of political opposition that structure people's political opinion. The adequacy of RCA to study political opinion structures has begun to be increasingly recognized by empirical studies on public opinion; in recent times, several investigations have applied it to describe affective polarization patterns (Brensing and Sotoudeh 2022).

RCA certainly represents an important step forward to robustly analyze recurrent patterns of attitudinal heterogeneity; however, it also carries several drawbacks when it is applied to data that exhibit two features that are dominant in the type of data typically available to study public opinion. The first is answer rank heterogeneity: public opinion questions present to respondents different numbers of categorical answers to choose from (typically, 2, 5 or 7). The second is bipolarity, a term that refers to questions that ask survey takers either to choose between answer points that express magnitudes of rejection or support to question statements (Ostrom, Betz, and Skowronski 1992; Malhotra and Krosnick 2007.¹). 70% of the public opinion items contained in the codebook of the most recent ANES time series cumulative data file ANES 2022 exhibit a bipolar data structure.

1. A more formal definition of bipolarity describes it as a form of questioning that offers a "linearly ordered set between elements considered as negative and positive categories" (Batyrshtin et al. 2017)

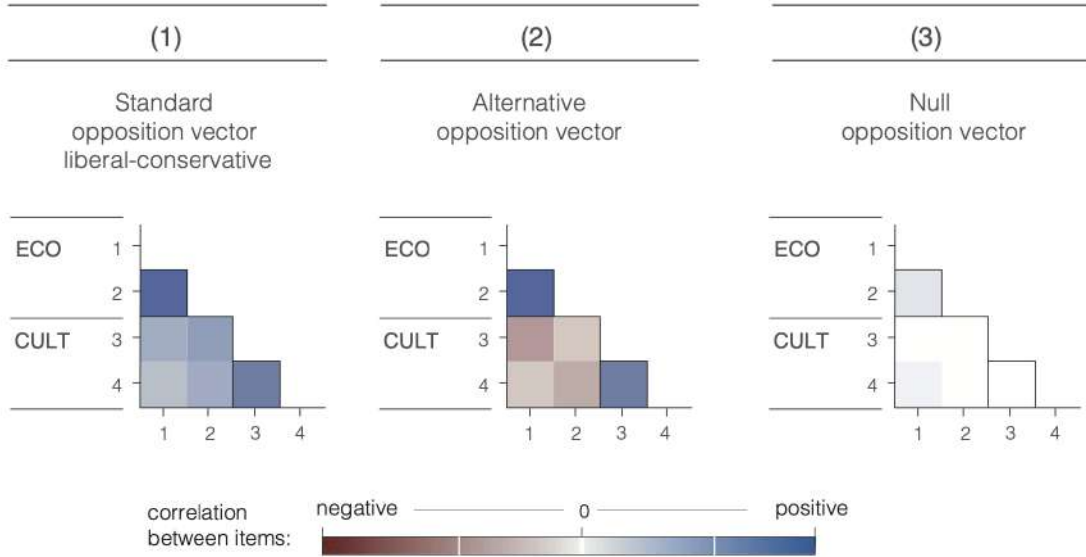


FIGURE 2. Expectations about political opposition vectors.

Notes: Graphs describe stylized expectations on patterns of opinion correlation for a simplified survey environment of 2 economic policy and 2 "cultural war" matters questions. The chart in the left depicts a standard left-right vector of political opposition where all opinions are significantly and positively correlated. The one in the middle displays an "alternative" vector where opinions on cultural positions and economic policy are negatively correlated. The chart on the left shows a hypothetical configuration of a "null" vector of political opposition where opinions across issues are not significantly related

We identify two inferential issues in RCA in environments characterized by answer rank heterogeneity and bipolarity.

One of these issues stems from the transformation of respondent answers into bounded variables with range $[0,1]$. When answer rank heterogeneity is present, this operation assigns different values to identical opinion expressions—and, conversely, identical values to dissimilar ones. For instance, looking at fFigure 3, in a 5-point Likert scale item the expression "agree" receives a value of 0.75, but in a 3-point question the same answer is recoded onto a value of 1. This inconsistency in valuation rolls over to inconsistent response similarity measurements, which induces in turn unreliability in the partition of the similarity graph into clusters representing VPOs. ²

2. Suppose, for instance, a situation where two respondent select (**Agree**, **Agree**) and (**Disagree**, **Agree**) as answers for two questions with 3 possible answers. Their answer would be recoded as (1,1) and (0,1), respectively; these numbers would yield a relationality score of 0 between them. By contrast, if these answers would have been choices out of 7-point questions, they would now valued as (5/6,5/6) and (1/6,5/6). These scores would produce a pairwise similarity score of (1/3).

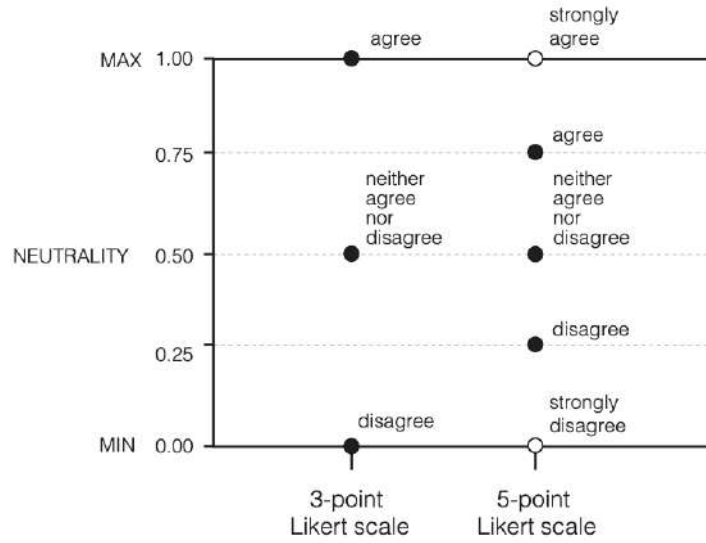


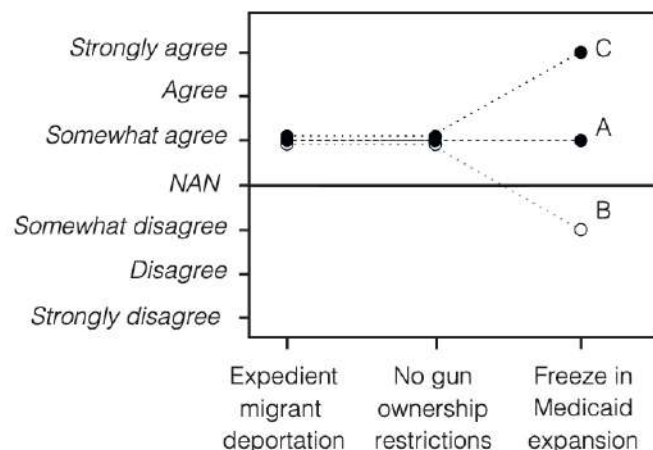
FIGURE 3. Mapping opinion answers into bounded cardinal variables induces reliability issues in the presence of answer point heterogeneity.

Notes: The diagram represents the numeric values that answer statements for 3- and 5-point Likert scale items would be recoded onto by RCA. It shows that these procedures assign different numeric values to same answer statements and identical values to different ones.

The second source of inferential drawbacks is the metric used by RCA to measure response pattern similarity between respondent pairs. This metric was introduced by Goldberg with the name of "relationality". Relationality averages pairwise scores of answer similarity that measure, for each pair of questions, how close the distance between answer choices picked by one respondent was relative to another (Goldberg 2011, 1406). These computations face inferential challenges on two grounds. First, they interpret recoded numerized answers as cardinal variables that can be compared in terms of magnitude even if they can only be accurately contrasted in terms of order. People choose between "slightly", simply, or "strongly" agreeing or disagreeing to a issue posed by a question according to thresholds of opinion intensity that vary across topics and individuals. Against this backdrop, examining response pattern similarity in terms of co-variations in answer distances are prone to calculate similarity scores that might incorrectly interpret people's positions as more distant (or closer) than what they are are.³

3. While Goldberg (2011, Appendix E) recognizes this issue, he argues that RCA—and, by extension, subsequent methods—can still be applied to analyze questions with 5 or more answer options because at this rank value ordinal categorical variables start approximating continuous ones. We believe this line of argumentation still carries two limitations. First, respondents' "sign first, then magnitude" answer selection strategy considerably reduce the number of answer options that respondents truly consider as answer choice possibilities. In a 7-point Likert item, for instance, this number is only 3. Second, statistical findings supporting that ordered variables can be analyzed as continuous

I. Opinions from respondents across issues



II. Relationality value for respondent pairs

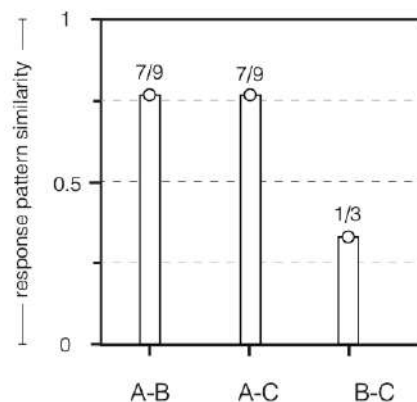


FIGURE 4. Political opinion answers in a hypothetical survey example.

Notes: nomenclature, black dots refer to respondents that rank as “standard ideologues” in Baldassarri and Goldberg’s nomenclature; white dots indicate those characterizable as “alternative ideologues” (see Baldassarri and Goldberg 2014).

More generally, understanding response pattern similarity in terms of co-shifts in answer distances is also problematic irrespective of the treatment of answers as cardinal variables because it leads to mistaken calculations of opinion similarity for bipolar answers. We discuss these problems using an illustrative example that analyzes resemblance in response patterns between individuals A, B, and C, who take part in a public opinion survey in the US. The survey asks them to choose among seven possible options (strongly Disagree, Disagree, Somewhat Disagree, Neither Agree nor Disagree, Somewhat Agree, Agree, and Strongly Agree) to react to the following statements: “Undocumented migrants should be summarily deported” (item 1), “Gun control should be free of background checks” (item 2), and “Medicaid should not be further expanded” (item 3). The left chart in figure 4 plots answer values for each respondent.

after a certain rank threshold (1987; 1981 have been produced only for correlation coefficients and for contexts where these variables follow a joint normal distribution (which might or not be the case for public opinion answers, as Goldberg asserts (Goldberg 2011, Appendix A). Similar results have yet to be produced for the specific types of operation that RCA procedures are based on. It is also worthwhile noting that, in practice, RCA has been used to analyze survey items with less than 5 answer points. Baldassarri and Goldberg 2014, analyze survey data where 46.5% of the questions have answer ranks lower than 5).

We draw from Baldassarri and Goldberg (2014) to assume in this example that the polity where respondents live has a standard left-right VPO that opposes overarching progressives who strongly reject expedient migrant deportation, easy access to guns, and freezes in public health care expansion with overarching conservatives” who strongly support all of these policies. In our example, individual A ranks as a member of the left-right divide; she can be characterized as a moderate conservative. Her answers are (Somewhat Agree, Somewhat Agree, Somewhat Agree). Person C also ranks as a standard moderate conservative, but she voices a stronger degree of support for the last question. The opinions she expresses are the following: (Somewhat Agree, Somewhat Agree, Strongly Agree).

The second VPO is inhabited by ”alternative ideologues” that exhibit positions on items 1 and 2 that positively related to each another but negatively related to item 3. In one pole of this vector lie people who support fast migrant deportation and no changes in gun control policies but object halting Medicaid expansion; in the other are those who reject expedient migrant expulsion and easy access to guns but oppose further expansions of Medicaid. Respondent B’s views make include here in this PVO. She is a moderate redistribute conservative whose set of answers is (Somewhat Agree, Somewhat Agree, Somewhat Disagree).

The right chart in figure 4 plots the scores of response pattern similarity that relationality gives to each respondent pair. We note that this measure assigns identical values to (A,B) and (A,C) even if only the latter rank as members of the same vector of political opposition. The results above show that relationality is prone to see people that carry different structures of political opinion as more similar than people who indeed share similar patterns of political opinion variation. Since relationality scores are the core input that RCA uses to assemble a similarity graph between respondents, this situation ends up compromising the validity of the response pattern clusters produced by it.

The locus of the validity problems described above lies in the fact that relationality is a distance-based measure of similarity that is indifferent to whether answers’ movements across rejection and support opinions are close or not with each other. This indifference does not adjust well with what we know of how people go about choosing answers to survey questions. The literature on survey item response indicates that, to answer bipolar questions, survey respondents use a “sign first, magnitude second” decision-making process: they first assess whether they locate themselves in a rejection or support position, and they choose the answer that best expresses the intensity of this rejection or support position *afterwards* (Ostrom, Betz, and Skowronski 1992). These sequential answer selection strategy indicates that an adequate metric of response pattern similarity

should prioritize capturing how close respondents’ answers change between support and rejection answers. Relationality underperforms by this criterion: it can’t keep track of how changes between rejection and support spaces co-vary between respondents. In our example, this is the reason why RCA assigns the same similarity value to respondents (A,B) and (A,C): respondent B’s answer about Medicaid expansion (**Strongly Agree**) is “two opinion points” higher than respondent A’s (**Somewhat Agree**), while respondent C’s answer (**Somewhat Disagree**) is two “opinion points” lower than respondent A’s. RCA treats these relative positions as equivalent even if they lead to substantive differences in how respondents feel about this issue.

SECTION 3. BIPOLAR CLASS ANALYSIS

We develop Bipolar Class Analysis (BCA) as a new method capable of overcoming RCA’s inferential limitations. BCA is explicitly designed to analyze data environments characterized by bipolarity and answer rank heterogeneity. It is premised on the characterization of answer sets from bipolar items as a union of a positive opinion semispace $\mathcal{P}$ and a negative opinion semispace $\mathcal{N}$ that share in common a neutrality point (typically in the form of a “Neither Agree nor Disagree” answer option).⁴ This option is conceptualized as being shared by $\mathcal{P}$ and $\mathcal{N}$ because it communicates a *neutral* position relative to a question statement.⁵

We offer a detailed formal description of BCA in Appendix A; the paragraphs below are focused on presenting it in an intuitive fashion. BCA adopts RCA’s operative algorithm but adopts two distinctive procedures at the first two steps of this sequence. In the first step, it does not demand answer variables to be mapped onto bounded cardinal variables; it requires only identifying a neutral point dividing responses into positive and negative semispaces. In the second steps, it replaces relationality by a *polar* measure of response pattern similarity $S_{polar}(u, v)$ that takes the following functional form:.

$$(3.1) \quad S_{polar}(u, v) = \frac{2}{Q(Q-1)} \sum_{k=1}^{Q-1} \sum_{l=k+1}^Q \omega(u_{kl}, v_{kl}), .$$

4. We consider that, even if it is not an explicit option, there exists a neutral position $\bar{n}_q \in [0, 1]$ for each question. The neutral position models the consensus of the group regarding questions q .

5. Strictly speaking, the opinion semispaces we distinguish refer to non-negative and non-positive opinion values; we use the labels “positive” and “negative” for ease of exposition. In the mathematical framework we propose, the neutral point may be explicit, as the questions of our motivating example, or implicit, as in binary questions that offer only favor or against answer options.

BCA’s polar similarity measure averages pairwise similarity scores, denoted by the term $\omega(u_{kl}, v_{kl})$. They act as a sign function that evaluate how close the answers that a respondent gives to a pair of questions shift across opinion semispaces relative to the answers from another one. These terms take the values of $-1, 0$, or 1 depending on how answer shifts vary between respondents. Figure 5 graphically presents the values that ω takes across all such possible shifts. In a nutshell, ω is -1 when a respondent’s answers to a pair of questions shift in a way that makes them be in opposite semispaces relative to the answers of another survey taker (panels I and II in Figure 5; 0 if answers to a pair of questions change semi spaces for one respondent but not for the other (panels III and VI); and 1 if one respondent’s answers to a pair of questions vary in such a way that makes them share the same semi-space than the answers from another one (panels IV, V, VII and VIII).

Returning to our working example in section 2, we can see that BCA’s polar similarity measure does a better work in accurately capturing likeness in response patterns. Looking at Figure ??, we can see that bipolar similarity measures correctly give larger similarity values to pairs of people located in similar VPOs (A, C) and smaller values to couples who belong to different ones (A, B; B, C).

The above analysis suggests that the polar similarity measure we introduce can more accurately capture likeness in pattern response among respondents, and ultimately, allow BCA to generate more accurate estimations of VPO structures. To generalize this result, we evaluate the BCA’s ability to accurately estimate vectors of political opposition relative to RCA through the analysis of 1,000 simulated survey datasets with 10 questions and a respondent population between 200 and 2,000. Answer distribution across negative and positive semispaces for these questions is set to vary at random so that the positive answers fluctuate to be between 35 and 65% of the answer total. Key for our purposes, datasets also vary in terms of the number of VPOs present in the respondent population (2, 3, or 4) and the number of respondents ascribed to each VPO (between 100 and 500). These attributes are known and predetermined. The structure of associations between answers for each VPO is also preset by design through a correlation matrix Σ^k , where correlations between positions are randomly set.

The data generation process (DGP) for the datasets described above follows three steps—a more in-depth technical discussion of it is available in Appendix B. The first specifies the number of respondents for each dataset. The second generates a matrix X^* of size $N \times Q$ of latent, mutually dependent opinion positions, and the third maps these latent positions onto an equally-sized matrix X of observed opinion answers. We highlight two characteristics in them that add realism to the DGP by capturing key features of the cognitive processes that individuals use to respond to

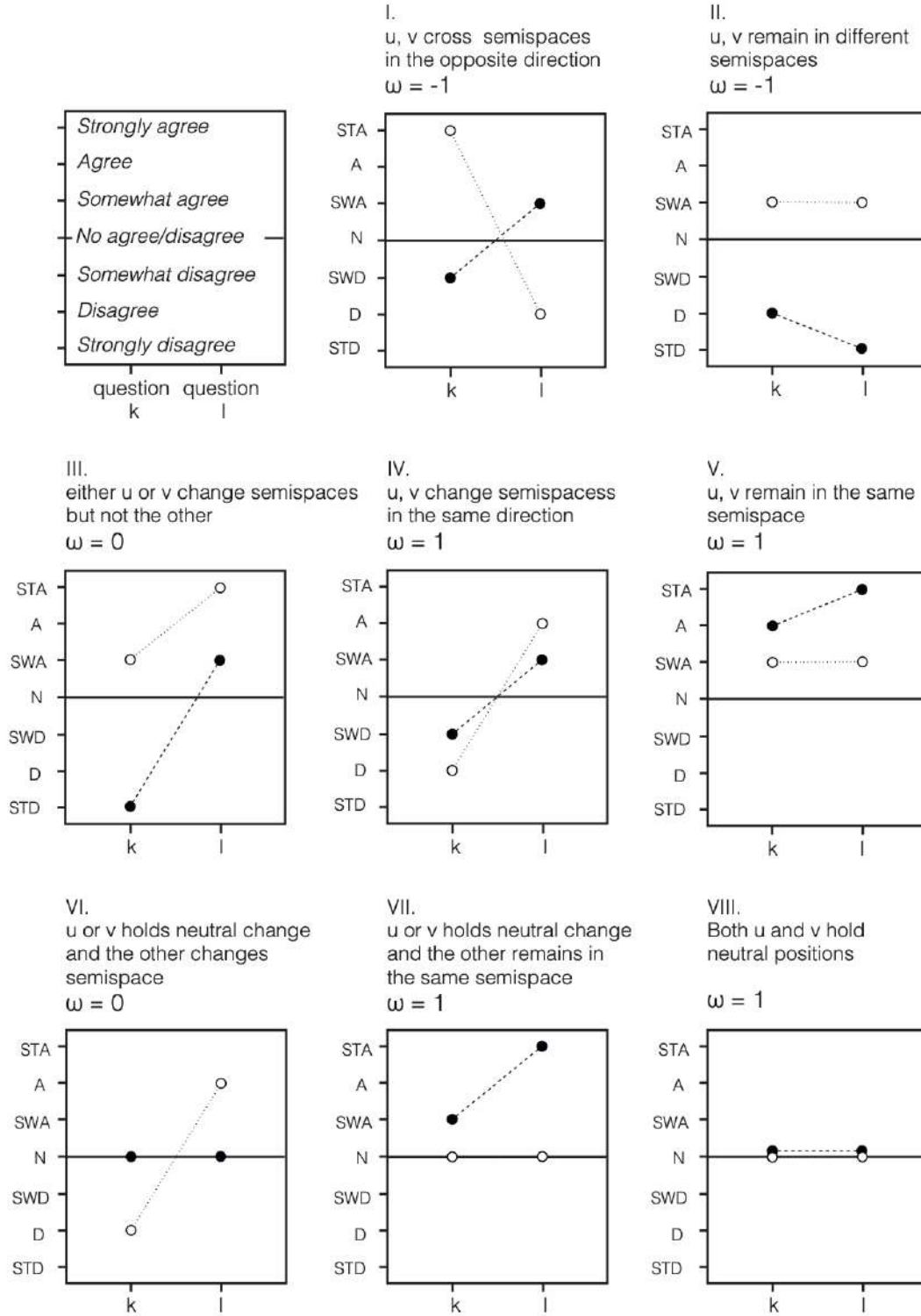


FIGURE 5. BCA: polarity function ω .

Notes: Black dots represent answers of respondent u to question items k and l ; white ones represent answers of a hypothetical respondent v for the same question.

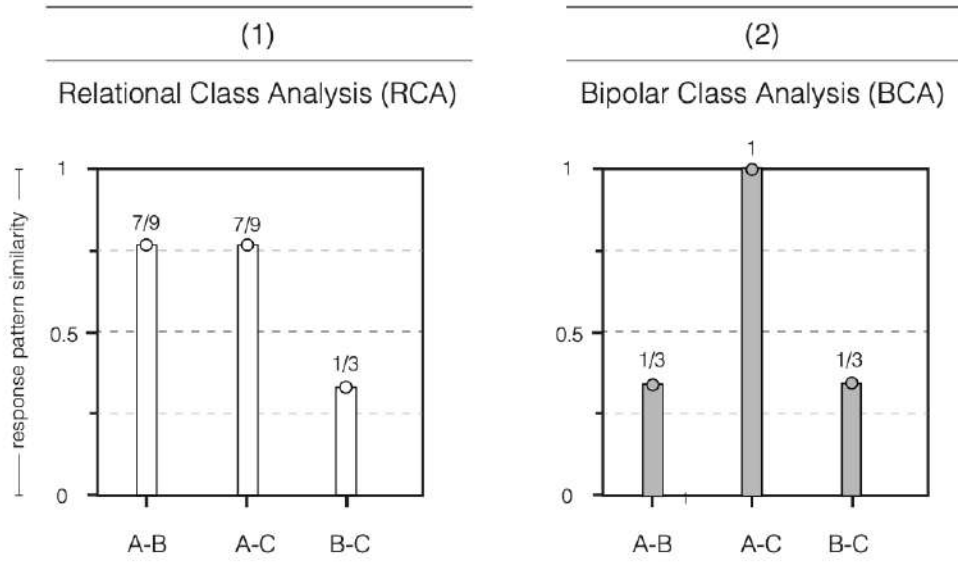


FIGURE 6. BCA measurements of response pattern similarity.

Notes: BCA’s polar similarity measures are better than RCA in capturing relationships of likeness and dissimilarity in the response patterns of A, B, C.

surveys. One is *ordered choice* (Greene and Hensher 2010: when prompted to answer a public opinion question, respondents establish a latent, unobserved *position* that is then used to select a specific *answer* among a finite set of ordered alternatives. In our simulations, the conversion of positions into answers is conducted by individual-specific mapping processes. The second is *supra-linear dependency*: respondents’ positions on survey issues are interconnected by specific constraint (i.e., dependency) structures that are more flexible than the linear forms of dependency that have been previously used to model attitudinal associations.

We compare BCA and RCA results across three parameters of performance. The first is *Partition accuracy* which is the rate at which a method rightly estimates the number of VPOs across datasets. The second is the *Scaled Normalized Mutual Information* (Scaled NMI). It is a variant of a measure of accuracy in observation classification across clusters known as Normalized Mutual Information (NMI). Values for this statistic range from 0 to 1; bigger numbers indicate a higher ability to assign a respondent to the VPO that his political opinions makes her a part of.⁶ The third and final

6. The NMI has been frequently used as a performance parameter in past research on RCA and related methods (Boutyline 2017; Sotoudeh and DiMaggio 2023), but several methodological research has warned that that this parameter produces upwardly biased estimations of classification accuracy when a large number of cluster classes are estimated (Amelio and Pizzuti 2015; Boutyline 2017, p. 388; Mahmoudi and Jemielniak 2024). Thus, based on the results from these studies and following their recommendations (Amelio and Pizzuti 2015), we use Scaled Normalized Mutual Information as an index of accuracy in classification. This number reports the Normalized Mutual Information

performance statistic we report is *correlational dissimilarity*. It assesses how proximate estimated constraint structures of VPOs are to the characteristics of the known dependency structures of "true" known VPOs in the data set. This index is discussed in detail in Appendix C. If the estimated and true correlation matrices are equal (i.e., they are identical), correlational dissimilarity takes the value of zero; lower correlational dissimilarity values indicate more accuracy in the estimation of dependency structures.

	(1) RCA	(2) BCA
Partition accuracy	0.046	0.410
Correlational dissimilarity	5.891	4.012
Mean Scaled NMI	0.136	0.214

TABLE 1. CA outperforms RCA across all measures of clustering accuracy.

Notes: Clustering accuracy performance across simulation experiments; mean values across experiment simulations are shown.

Table 1 shows the results of the simulation analysis. We see that RCA accurately predicts the number of underlying political opposition vectors in the data only 4.6% of the time. BCA’s rate of accuracy is much higher, sitting at 41.0%. BCA also ranks considerably better in producing correlational structures closer to the true ones. The mean correlational dissimilarity score is 4.012, roughly a third less than RCA’s (5.891). Results also cast RCA as a low performer in its ability to correctly estimate respondent’s VPO membership. The scaled NMI value for this method is 0.136. BCA’s is more than a half larger, sitting at 0.214 (CCA). Overall, these results present BCA_1 as a method that is more accurate in estimating the number correlation structure of VPOs, and respondent’s VPO membership.

SECTION 4. RESEARCH DESIGN

We use answers to public opinion questions from the 2020 ANES time-series cumulative dataset to investigate how VPOs developed across 6 time points between 1988 and 2020. We conduct time-series analyses of different historical width that apply BCA to partition respondents by political

in a scaled form that adjusts NMI values downward when the estimated number of groups significantly exceeds the correct number

response similarity for each each of the time-points they include using answers to identical sets of variables. This strategy allows addressing temporal comparability issues related to compositional heterogeneity. They also allow exploring the stability of results when BCA is used to analyze different sets of answer variables.

Variable selection. ANES surveys include public opinion questions in a very irregular fashion. None of these questions has been included in all of the 17 ANES waves with information on people’s political opinions between 1980 and 2020;⁷ they were included on average in only 8. The inclusion of public items is also highly discontinuous: the average span of consecutive inclusion in ANES waves among them is only 3.05. The irregularity in the inclusion of public opinion items dovetails also with large differences in the number of this type of questions across survey waves. On average each wave included 72 questions with a standard deviation of 21. These figures indicate that location of public opinion items across ANES is scarce, sparse and unsystematic. This is a challenging data environment that makes historical analyses on patterns of political opinion heterogeneity highly exposed to comparability issues related to compositional heterogeneity patterns of political opinion. This problem, which has gone unrecognized in previous research, prompts up to give variable selection strategy more careful consideration than the one it has typically received.

We address compositional heterogeneity in our analysis by following a 4-step variable selection procedure.

The first step identified and mapped the inclusion over time of all ANES survey items measuring opinions on political issues between 1980 and 2020.⁸ We identified 95 such questions. They asked opinions on seven different areas: economic policy; civil rights and black disadvantage (henceforth “civil rights”); foreign policy; law and order; migration; environmental policy; and topics related to women’s role in society, LGTBQ+ discrimination, and support traditionalism that tend to fall under the rubric of “culture war issues” in contemporary American public discourse; law and order; migration; and environmental policy.

The second step concerned time point selection. We were interested in identifying a set of time points whose analysis increased the likelihood of counting with a high number of public opinion variables fully included across them while at the same time sitting at temporal points that allowed assessing large-span historical evolution of VPOs. With this in mind, we identified the ANES wave conducted in a presidential year that included the highest number of public opinion questions for

7. the ANES trendfile does not report public opinion questions for mid-term waves after 2006

8. We understand an item to be gathering opinions on political issues if they ask about people’s stances on policy issues, preferences of desirable states (for example, equality or a traditional way of living), or assessments of political interest groups directly associated with specific policies (for instance, unions or feminist activists)

each decade starting in the 1980s. Such years were 1988, 1992, and 2008. We selected these years as our target time points. We also selected 2012, 2016 and 2020 since they were the most recently conducted ANES surveys and also because they have density of public opinion questions (89, 81, and 72) comparable to those with the highest number in the preceding decades. Thus, in total, we selected six years as our time points of interest.

The third step identified the issues targeted by the public opinion questions that were included in the ANES waves that took place in the years of interest. Several questions targeted the same issue—Consider, for instance, items VCF9015, VCF9016, and VCF9017. The first asked how problematic it was “that not everyone has equal chance”; the second, agreement or disagreement that it “was not a big problem if some have more chance in life”, and the third, whether a respondent thought that people “should worry less about how equal people are” (ANES 2022, , 500-502). These questions were coded as gathering opinion on the issue of inequality’s potential social harm. Overall, we identified that survey items of interest gathered opinions on 58 different issues.

In the fourth and last step, we identified sets of issues that were included as public opinion issues in every data point of interest since 1988, 1992, 2008, and 2012. These sets are listed in Table 7. We will refer to each of these four sets of issues by the year at which they were fully included (for instance, the “1988 set” refers to the set of 19 topics that were fully included from 1988 to 2020; the “2012 set” refers to the 29 themes for whom public opinion information was fully available from 2012 to 2020. Variables that gather opinion information on each issue were chosen as our variable set of interest.

Analytic strategy. We adopt a “time-series comparison” framework to analyze the data. We applied BCA to analyze opinions for each set of issues across the years when they were consistently included. This led to the mapping of VPOs for 25 different time point/variable set combinations. Using these mappings as inputs, we then assessed how the VPO structures that were estimated for each set of issues developed over time. Since there is a trade-off between topic and temporal range across these issue-set time series, in the last step we evaluated how close the development of estimated VPOs over time was among these time series. We sought to address compositional heterogeneity and maximize temporal comparability of results by focusing first how results within each issue sets evolved over time; we then considered how robust temporal results were to different issue specifications by comparing results across issue sets.

We analyze the historical development of estimated VPOs to explore three substantive questions. A first question concerns observed diversity in VPOs. Is there evidence supporting that people’s political oppositions are organized around a standard lib-con, null, or alternative vector of political

Area	Issue	
Issue set 1988: topics consistently surveyed across 6 selected ANES data points, 1988 - 2020		
Economy	1	Gov't capacity to efficiently handle resource
	2	Trade unions
	3	Big corporations
	4	Public health
	5	Guaranteed income and employment
	6	Tradeoff between public spending and extending public services
	7	Public education
	8	Economic protectionism
Equality	9	Relevance of inequality as a social problem
Civil Rights	10	Affirmative action
Culture	11	Feminist activism
	12	Support for abortion
	13	Gay antidiscrimination legislation
	14	Christian fundamentalism
	15	Moral readjustments according to cultural changes
	16	Promotion of traditional values
Law and order	17	Death penalty
Migration	18	Undocumented migrants
Environment	19	Funding environmental policies
Issue set 1992: topics consistently surveyed across 5 selected ANES data points, 1992 - 2020		
	1-19 (Issue set 1988)	
Economy	20	Public investment in poverty relief
Migration	21	Increase in the volume of new legal migrants
Law and order	22	Police forces
Issue set 2008: topics consistently surveyed across 4 selected ANES data points, 2008 - 2020		
	1-22 (Issue set 1992)	
Migration	23	Migration and unemployment
Law and order	24	Gun control
Civil Rights	25	Influence of black people in society
Issue set 2012: topics consistently surveyed across 3 selected ANES data points, 2012 - 2020		
	1-25 (Issue set 2008)	
Economy	26	Public investment in social security ¹⁸
	27	Breadth of government activity
Civil rights	28	Influence of slavery and past discrimination in current black disadvantage
	29	Fairness in black share in society

opposition? We assess this issue by analyzing *constraint rates* (CRs) for estimated VPOs, which we define as the percentage of opinion correlations between issues that are statistically significant. We understand a *null VPO* as one with a constraint rate smaller than 0.2. In addition, we characterize a standard "lib-con" as one where the share of positive significant correlations between issue opinions—henceforth: *positive constraint rate*, or PCR—is larger than 0.5. We characterize vectors that are neither lib-con nor null as "alternative" ones.

A second question concerns temporal development: how have our VPO categories of analytical interest—lib-cons, nulls, or alternatives—changed over time? Has each become more marginal or more dominant? We study this question by exploring changes in general and positive constraint rates, on one hand, and shifts in their relative size, on the other. These variations provide information on whether lib-con, null or alternative VPOs have become "stronger" via *intensification*, *growth* or *diffusion*. For the case of a lib-con lines of political opposition, strengthening through intensification would mean increases in positive *constraint rates* among estimated vectors; through growth, to increases in the *relative population* of vectors ranking as lib-con; and through *diffusion*, to increases in the *number of estimated vectors* ranking as lib-con. ⁹.

SECTION 5. RESULTS

We estimate vectors of political opposition for the ANES data points of interest using the 1988 issue set. This is the largest time-span analyses we run. Figure 8 plots correlation matrices for each estimated VPO. In each year VPOs are ordered from left to right by positive constraint rate; the first for each year is thus the one ranking highest in this measure. The stars to the left of Y-axis indicate the strength to which a vector ranks approximates a standard lib-con opinion pattern. A larger number of stars indicates stronger lib-con resemblance. Lack of stars indicates that rate positive constraint is smaller than 0.5.

We begin analyzing BCA results by centering on the constraint qualities of vectors that rank highest in PCR. The PCR value for such vector in 1988 is 0.47, just below the threshold to classify as lib-con. We can thus describe this vector as an "alternative" configuration of opinion patterns. ¹⁰. For later years, all the vectors with the highest PCR classify as lib-con. PCR scores across them increase consistently between 1992 (PCR = 0.73) up until 2012 and 2016 (PCR = 0.87 and 0.90, respectively). By contrast, we note a descent of this parameter to levels below 1992 in 2020 (PCR = 0.69).

9. we call analogue process of weakening *moderation*, *shrinking*, and *retraction*

10. the most important characteristic of this alternative vector are a sparse correlation between economic and culture war attitudes, and negative correlations between the latter and economic protectionism

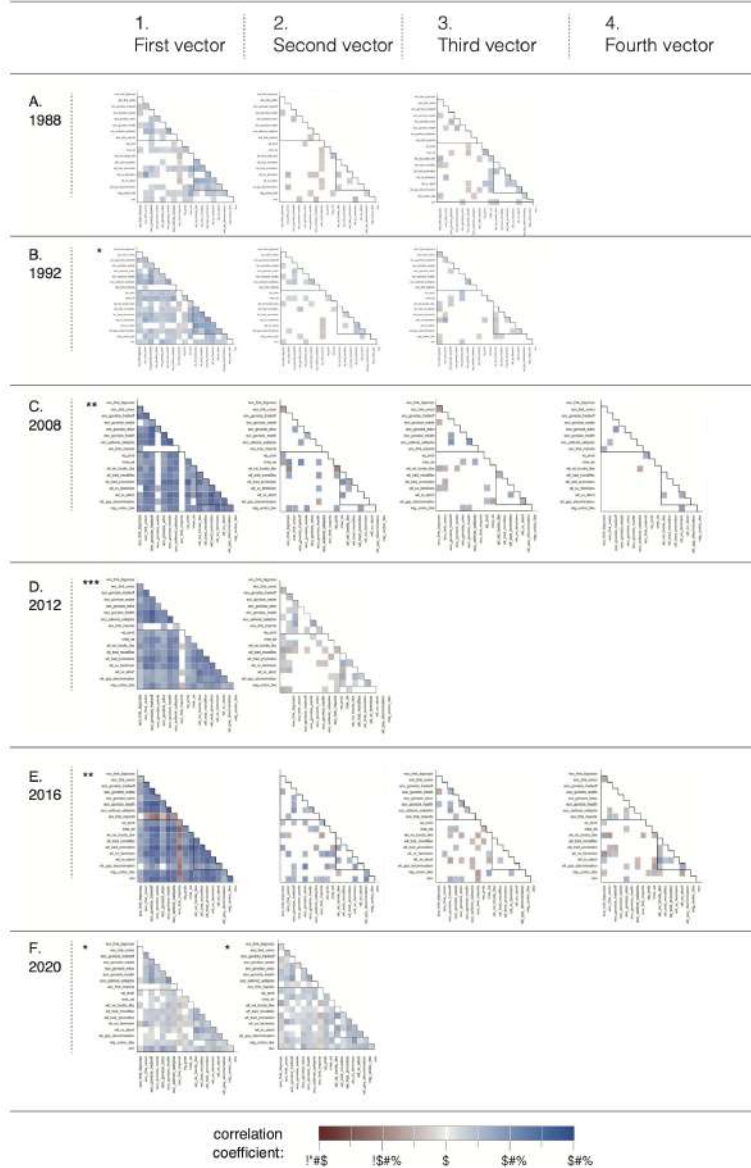


FIGURE 8. between 1988 and 2020, lib-con opinion structures have strengthened first via constraint intensification, and most recently, via dispersion onto additional VPOS.

Notes: The chart shows correlation matrices for estimated vectors of political opposition. Vectors are organized according to positive constraint rates (PCRs). Stars at the left of the x axis distinguish vectors that rank as lib-con opinion structures. $PCR_{l.50}$: *; $PCR_{l.75}$: **, $PCR_{l.90}$: ***

In all years, BCA also estimates between one and three vectors additional to the ones we described in the paragraph above. (Interestingly, we observe the largest density of such vectors in years that are recurrently described as highly transformative by virtue of having seen the election of Barack Obama and Donald Trump to the American presidency). All of these vectors rank as null: their

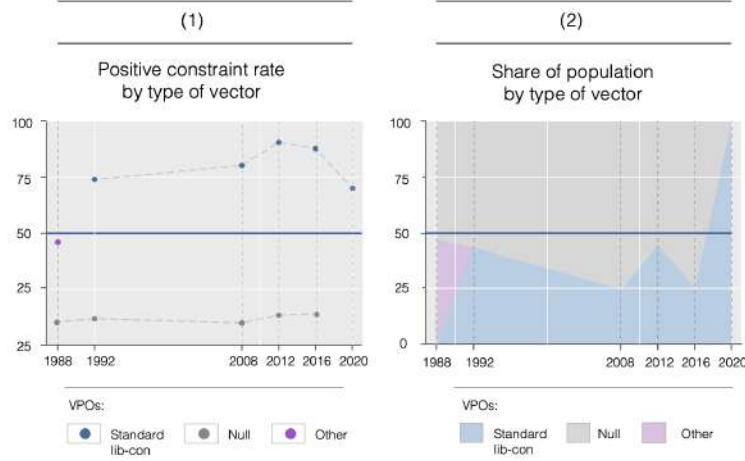


FIGURE 9. While constraint in standard lib-con axes has tended to go upwards since 1980s, this lib-con encompass less than half of the respondent population until 2020.

Notes: The left panel plots constraint rates for null and positive constraint rates for lib-con VPOs. Weighted averages are plotted for years when more than one vector ranks as part of these categories. The the plot in the right plots the relative population across vector categories.

constraint rate is consistently lower than 0.2. The only exception, and a significant one, for that matter, is 2020. The single additional vector for this year exhibits has a high PCR (0.69) and thus classifies as a lib-con vector.

In graph 9, the left chart visualizes the temporal location of lib-con, alternative, and null VPOs. We see that throughout most of time points under consideration, Americans tended to organize their political opinions mainly around lib-con and null structures. There are two significant exceptions: 1988, which lacks a lib-con VPO, and 2020, when all estimated VPOs ranked as part of this category. The chart also allows to see the evolution of constrain rates for null VPOs and trends in PCRs for those ranking as lib-con. We can see than, on average, the constraint rate of null vectors remained stationary around the 0.10 value. Standard lib-con vectors, on the other hand, exhibit show a non linear relationship from the the first time they appeared up until 2020.

The right chart in graph 9 plots the relative size of VPO types, measured as the share of respondents gathered by each. We can see that, up until 2020, the number of respondents ranking as members of the lib-con vector never attain a majority. This chart reveals that the observed intensification of positive constraint in lib-con vectors does not translate into size growths. We note in fact that moments when PCR is the highest are also the ones where lib-con vectors are most scarcely populated. In 2008, the lib-con vector exhibited a constraint rate beyond the 0.90

mark, but it gathered gathered by a quarter of the sample population. Indeed the majority of the population samples rank as members of null vectors for all time points except 2020, when all the sample ranks as members of lib-con vectors. Thus, in the most recent data point available for analysis, lib-con structures of political opinion appear to have become demographically dominant at the expense of a strong positive constraint.

One may ask how much the results obtained above are sensitive to changes in the array of issues selected for analysis. To assess this, we also examine BCA results for the three sets of variables that were repeatedly asked since 1992, 2008, and 2012, respectively. We report the results for the 2012 set of variables: which allows inspecting whether VPO results for 2012, 2016 and 2020 were similar when drawn from an expanded list of topics.¹¹

Figure 10 plots the correlation matrixes for the VPOs that were estimated using the 2012 set of issues. We can see the number of estimation VPOs shrinks when we used this expanded set of topics. BCA procedures estimate two VPOs of roughly equal population across the 2010s. One consistently ranks as a lib-con vector. Looking at these vectors, PCR steadily grows between 2012 and 2016 from 0.63 to 0.77; four years later, this indicator had a similar value (0.75). In 2020, both of the two estimated vectors rank as lib-con. The years 2012 and 2016, on the other hand, include "alternative" (rather than null) political opposition vectors.¹² This is the main difference observed between the analysis based on the 2012 issue set and the one drawing from the 1988 list of issues. This result suggests variables related to law and order, migration and the environment act as a source of further constraint in the organization of political opinions.

As a preliminary general, we highlight that the results offer strong indication that lib-con political opposition structures have indeed become stronger over time via two stages. We can characterize the first one, spanning the period between 1988 and 2016, as a "specialization" pattern. In this time range, lib-con vectors become increasingly constrained but fail to expand to the general population. In the second stage, which is much more recent (2016-2020), lib-con vectors underwent a process of "generalization:" while constraint across opinions in these type of vectors remained stable in 2020, they managed to encompass the full set of the survey population. The results we report so far in this analysis are mixed relative to the persistence of null and alternative VPOs: results on these types of opinion structures appear to be sensitive to the opinion issues used for clustering respondent along response pattern similarity.

11. Results for the 1992 and 2008 issue sets, which are not shown, yield similar results to the ones reported below. They are available upon request

12. these alternative forms of political opposition are characterized by a negative correlation between positions on the provision of public funds for poverty relief and economic protectionism, on one hand with the rest opinions on civil rights, culture war items, and law and order.

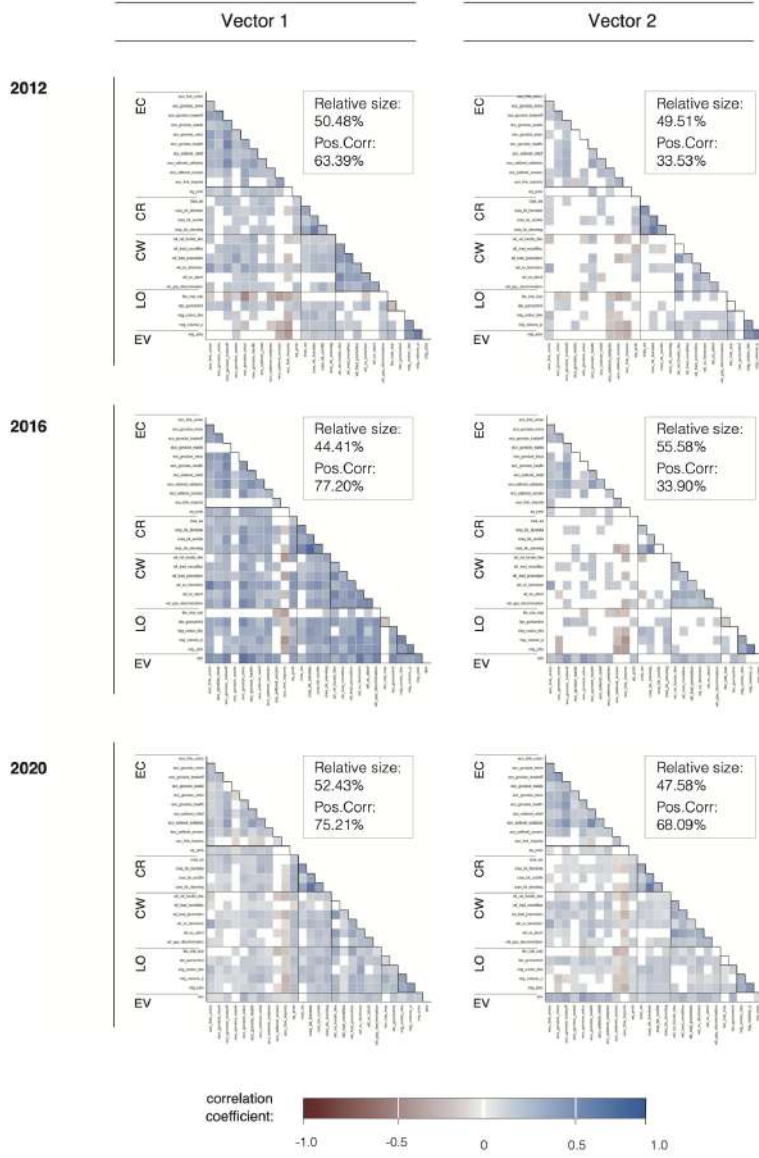


FIGURE 10. The standard liberal-conservative VPO has become increasingly salient in the US since the 2012.

Notes: The chart shows correlation structures for each estimated VPO in 2012, 2016, 2020. Items are grouped by policy areas. EC: economy; CR: civil rights; CW: "culture war" items; LO: law and order; EV: environment.

APPENDIX A. TECHNICAL EXPOSITION OF BIPOLAR CLASS ANALYSIS (BCA)

Scope and Data Structure

BCA generates estimations of political opposition vectors from data on $i = 2, 3, \dots, n$ survey respondents answering $q = 2, 3, \dots, Q$ public opinion survey items with a set $\mathcal{A}_j$ of bipolar answer

points. The number of answer points available can be two, as in a binary *favor–against* question with an answer set $\mathcal{A}_q = \{\text{in favor, against}\}$, or larger than two, as in the case of a standard 7-point Likert item. When considering a set of individual responses to a poll of Q questions, we can express the set $\mathcal{A}$ of answer options that BCA uses as the Cartesian product $\mathcal{A} = \mathcal{A}_1 \times \cdots \times \mathcal{A}_Q$. We call an element in $\mathcal{A}$ an *answer set*. We call $\mathcal{A}$ the *answer space*.

We note that the answer set $\mathcal{A}_\Pi$ exhibits two key mathematical properties:

Neutrality. $\mathcal{A}_\Pi$ can be characterized as a union of two distinct non-empty subsets: $\mathcal{A}_\Pi = \mathcal{N} \cup \mathcal{P}$ and $\mathcal{P}$. We call $\mathcal{N}$ the *negative* and $\mathcal{P}$ the *positive* opinion semispace, respectively. We identify $\mathcal{N}$ and $\mathcal{P}$ as subsets that are either disjoint or share a unique element in common. In the first case, $\mathcal{N}$ and $\mathcal{P}$ do not share any element in common. This occurs, for instance, in binary favor-against questions. In these cases, the neutrality element n is considered to be implicit.¹³ In the latter, $\mathcal{N}$ and $\mathcal{P}$ include an answer point that expresses a neutral position to a question statement. We call this option the *neutrality element*, denoted by n . In a 7-point Likert question, the answer point “neither agree nor disagree” is the neutrality element. In this case, $\mathcal{P} = \{\text{neither agree nor disagree, slightly agree, agree, strongly agree}\}$ and $\mathcal{N} = \{\text{strongly disagree, disagree, slightly disagree, neither agree nor disagree}\}$

Total order. The answer points in $\mathcal{A}_q$ are a *totally ordered set*, which we express as a binary relation $\leq_q$ that satisfies the following conditions:

- *reflexivity:* $a \leq_q a$ for all $a \in \mathcal{A}_q$,
- *transitivity:* given $a, b, c \in \mathcal{A}_q$, if $a \leq_q b$ and $b \leq_q c$, then $a \leq_q c$,
- *anti-symmetry:* given $a, b \in \mathcal{A}_q$ with $a \neq b$, if $a \leq_q b$, then $b \not\leq_q a$,
- *strong connectedness:* for any $a, b \in \mathcal{A}_q$, either $a \leq_q b$ or $b \leq_q a$;

It is straightforward to verify that bipolar items—for example, the set of 7 answers for *strongly agree–strongly disagree* questions—support a natural total order given by

$$(A.1) \quad \text{strongly disagree} \leq \text{disagree} \leq \text{slightly disagree} \leq \text{neither agree nor disagree} \\ \leq \text{slightly agree} \leq \text{agree} \leq \text{strongly agree}.$$

Based on the neutrality and total ordered conditions of bipolar survey items, we can express an answer set $\mathcal{A}_q$ with an explicit neutrality element n_q as the union $\mathcal{A}_q = \mathcal{P}_q \cup \mathcal{N}_q$ with

13. Decomposing an answer set $\mathcal{A}_q$ into $\mathcal{N}$ and $\mathcal{P}$ semispaces accommodates the possibility of heterogeneous number of elements among them

$$(A.2) \quad \mathcal{P}_q = \{a \in \mathcal{A}_q : n_q \leq a\}$$

and

$$(A.3) \quad \mathcal{N}_q = \{a \in \mathcal{A}_q : a \leq n_q\}.$$

Response similarity measurement in RCA

Consider a set $\mathcal{A}$ of responses from 1, 2, ..., n individuals to $Q \geq 1$ questions. RCA clusters respondents using relationality (R) as a measure of response similarity between two individuals. This method requires each answer set $\mathcal{A}_q$ in $\mathcal{A}$ to have a one-to-one correspondence with a strictly increasing finite sequence of equidistant numbers in the unit interval $[0, 1]$ (Goldberg 2011, 1406). Answer points in $\mathcal{A}_q$ are thus taken as a tuple of N numbers between zero and one, which allows the arithmetical manipulation of observed answers.

Let $u = (u_1, \dots, u_Q)$ and $v = (v_1, \dots, v_Q)$ be sets of observable answers in $\mathcal{A}$ picked by individuals u and v . We also define $u_{kl} = (u_k, u_l)$, and $v_{kl} = (v_k, v_l)$ for $1 \leq k < l \leq Q$. In this setting, R is a function $R_G: \mathcal{A} \times \mathcal{A} \rightarrow \mathbb{R}$ given by

$$(A.4) \quad R(u, v) = \frac{2}{Q(Q-1)} \sum_{k=1}^{Q-1} \sum_{l=k+1}^Q \lambda(u_{kl}, v_{kl}) \delta(u_{kl}, v_{kl}),$$

where

$$(A.5) \quad \lambda: (\mathcal{A}_k \times \mathcal{A}_l) \times (\mathcal{A}_k \times \mathcal{A}_l) \rightarrow \{-1, 1\}$$

$$(u_{kl}, v_{kl}) \mapsto \lambda(u_{kl}, v_{kl}) = \begin{cases} 1, & \text{if } (u_k - u_l)(v_k - v_l) \geq 0 \\ -1, & \text{otherwise,} \end{cases}$$

and

$$(A.6) \quad \delta: (\mathcal{A}_k \times \mathcal{A}_l) \times (\mathcal{A}_k \times \mathcal{A}_l) \rightarrow [0, 1]$$

$$(u_{kl}, v_{kl}) \mapsto \delta(u_{kl}, v_{kl}) = 1 - ||u_k - u_l| - |v_k - v_l||.$$

In equation (A.5), the λ term can be substantively interpreted as a “direction function” that measures convergence or divergence in the direction of change among questions between respondents. The δ term, on the other hand, represents a “distance” measurement that captures convergence or divergence in the distance of opinion movements among questions.

Measurement of response similarity in BCA

BCA uses a polar similarity measure to gauge proximity in response patterns between respondent pairs. This similarity measure shares with RCA’s relationality a paired-comparison strategy to measure response similarity. But BCA’s polar similarity term uses an orientation term ω instead of the direction function λ and takes out the distance term δ . We discuss the orientation term ω in a setting where $u = (u_1, \dots, u_Q)$ and $v = (v_1, \dots, v_Q)$ are series of observed answers in $\mathcal{A}$ that were picked up by individuals u and v , and for pairs of answers $u_{kl} = (u_k, u_l)$ for $1 \leq k, l \leq Q$ and $k \neq l$.

The orientation function ω measures whether observed answers from respondents converge or diverge in terms of how they switch between opinion semispaces—that is, how they move relative to the neutrality element. There are four different possibilities of displacement between opinion semispaces. We formalize these scenarios as follows by introducing four *movement functions* as follows:

$$(A.7) \quad M_{++}(u_{kl}) = \begin{cases} 1, & \text{if } u_k \in \mathcal{P}_k, u_l \in \mathcal{P}_l \\ 0, & \text{otherwise;} \end{cases}$$

$$(A.8) \quad M_{--}(u_{kl}) = \begin{cases} 1, & \text{if } u_k \in \mathcal{N}_k, u_l \in \mathcal{N}_l \\ 0, & \text{otherwise;} \end{cases}$$

$$(A.9) \quad M_{-+}(u_{kl}) = \begin{cases} 1, & \text{if } u_k \in \mathcal{N}_k^\circ, u_l \in \mathcal{P}_l^\circ \\ 0, & \text{otherwise;} \end{cases}$$

$$(A.10) \quad M_{+-}(u_{kl}) = \begin{cases} 1, & \text{if } u_k \in \mathcal{P}_k^\circ, u_l \in \mathcal{N}_l^\circ \\ 0, & \text{otherwise.} \end{cases}$$

Note that, for any answer pairs u_{kl} , at least one of the functions listed above must take the value of 1. Here, $\mathcal{N}^\circ$ consists of all the answer points in $\mathcal{N}$ except the neutrality element. More formally,

$\mathcal{N}^\circ = \mathcal{N} \cap \mathcal{P}^c$, where $\mathcal{P}^c$ denotes the complement of $\mathcal{P}$. Similarly, $\mathcal{P}^\circ = \mathcal{P} \cap \mathcal{N}^c$. We call $\mathcal{N}^\circ$ and $\mathcal{P}^\circ$ the *strictly negative* and *strictly positive* opinion semispaces, respectively.

We next define a *joint movement function* between pairs of observed answers. For ease of notation, we will write $n_{kl} = (n_k, n_l)$, where n_i is the neutrality element in $\mathcal{A}_i$. The *joint movement function* $f_{A,B}$ assigns to the pair (u_{kl}, v_{kl}) the value 0 or 1 and is given by

$$(A.11) \quad \begin{aligned} f_{A,B}: (\mathcal{A}_k \times \mathcal{A}_l) \times (\mathcal{A}_k \times \mathcal{A}_l) &\rightarrow \{0, 1\} \\ (u_{kl}, v_{kl}) &\mapsto M_A(u_{kl})M_B(v_{kl}), \end{aligned}$$

where $A, B \in \{++, --, +-, -+\}$.

BCA's *polarity function* ω gives the value 0, -1, or 1 to the response pair (u_{kl}, v_{kl}) . In the main body of the paper, Figure 5 provides a graphic illustration of the values that the joint movement function returns across opinion movement scenarios. These values are formally assigned as follows:

If $u_{kl}, v_{kl} \neq n_{kl}$:

$$(A.12) \quad \omega(u_{kl}, v_{kl}) = \begin{cases} 1, & \text{if } f_{A,A}(u_{kl}, v_{kl}) = 1 \text{ for at least one } A; \\ -1, & \text{if } f_{A,A^*}(u_{kl}, v_{kl}) = 1 \text{ for at least one } A; \\ 0, & \text{otherwise.} \end{cases}$$

Here, $*$ performs the following transformations:

$$(A.13) \quad (++)^* = --, (--)^* = ++, (-+)^* = +-, \text{ and } (+-)^* = -+;$$

If either $u_{kl}, v_{kl} = n_{kl}$:

$$(A.14) \quad \omega(n_{kl}, v_{kl}) = \begin{cases} 1, & \text{if } M_{++}(v_{kl}) = 1; \text{ or } M_{--}(v_{kl}) = 1; \\ 0; & \text{otherwise;} \end{cases}$$

$$(A.15) \quad \omega(u_{kl}, n_{kl}) = \begin{cases} 1, & \text{if } M_{++}(u_{kl}) = 1 \text{ or } M_{--}(u_{kl}) = 1; \\ 0; & \text{otherwise;} \end{cases}$$

and

$$(A.16) \quad \omega(n_{kl}, n_{kl}) = 1.$$

Summary of BCA

Summarizing, the polar measure of response similarity that BCA uses is the absolute value of $S_{polar}(u, v)$,

$$(A.17) \quad S_{polar}(u, v) = \frac{2}{Q(Q-1)} \sum_{k=1}^{Q-1} \sum_{l=k+1}^Q s^P(k, l)$$

$$(A.18) \quad s^P(k, l) = \omega(u_{kl}, v_{kl})$$

where ω is the polarity function defined in [\(A.12\)](#).

APPENDIX B. SIMULATION ANALYSES

Synthetic data modeling

We assess BCA’s performance relative to other CEMs using a large number of simulated datasets in which N respondents select answers to Q bipolar questions. As noted in section 3, the modeling strategy we use to generate these datasets has two distinctive features: ordered choice, on one hand, and a supra-linear dependency structure, on the other. We detail both characteristics in the paragraphs below.

Ordered Choice. We use a latent variable approach Greene and Hensher 2010 to model ordered choice in the way respondents pick their answers. The modeling process establishes a latent variable $x_{iq}^* \in \mathbb{R}$ that represents the accurate position for each respondent i on each issue q . The respondent then partitions $\mathbb{R}$ into Q intervals and chooses her answer $x_{iq} \in \mathcal{A}_q$ depending on the position interval where x_{iq}^* falls.

We formalize this procedure in the following way. For each respondent i and question q there are $H_q - 1$ thresholds $(\tau_{iq\eta})_{\eta=1}^{H_q-1}$, with $\tau_{iq1} < \tau_{iq2} < \dots < \tau_{iq,H_q-1}$ that partition $\mathbb{R}$ into H_q intervals. The respondent chooses x_{iq} depending on the interval where x_{iq}^* falls:

$$(B.1) \quad x_{iq} = a_\eta \iff x_{iq}^* \in (\tau_{iq,\eta-1}, \tau_{iq\eta}) \text{ for } \eta = 1, \dots, H_q.$$

In the equation above, we set $\tau_{iq0} = -\infty$ and $\tau_{iqH_q} = \infty$ for every i and q .¹⁴

Supra-linear dependency. We model political opposition vectors as specific dependence structures on the vector of latent positions $X_i^* = (x_{i1}^*, \dots, x_{iQ}^*)$. In other words, the mathematical equivalent to the sentence “respondent i belongs to a certain political opposition vectors” is “the answer variables in X_i^* , which show i ’s positions on each issue, follow a certain joint distribution”.

We apply the notion of copulas within probability theory to model dependence between the latent variables in X_i^* . Copulas are “a way of studying scale-free measures of dependence” Fisher 1997. A copula is a joint distribution on the unit cube $[0, 1]^Q$ with uniform marginal distributions that captures scale-free information about the dependence structure of random vectors Nelsen 2006. Based on *Sklar’s Theorem* (Th. 2.10.9 in Nelsen 2006), and letting F_q denote the marginal distribution of each variable x_{iq}^* , $q = 1, \dots, Q$, there exists a unique copula C such that

14. Also, since we will consider absolutely continuous distributions for both the latent variables and the thresholds, we note that $P(x_{iq}^* = \tau_{iq\eta}) = 0$. Thus, the fact that we consider open intervals is unimportant.

$$(B.2) \quad F(x_1, \dots, x_Q) = C(F_1(x_1), \dots, F_Q(x_Q)).$$

The equation above shows that the joint distribution of X_i^* can be decomposed into the marginal distributions F_q of answer variables and the copula C , which models the dependence structure on the joint distribution.

Data Generation Process

Modeling positions. Respondents' *positions* on questions are modeled as continuous random variables that range from 0 (full disagreement) to 1 (full agreement). Position values across questions are set according to a concrete dependency structure among them defined by the multivariate distribution of positions $F(x_1^*, x_2^*, \dots, x_Q^*)$.

Sklar's theorem allows decomposing the joint distribution of X_i^* into the marginal distributions F_q of answer variables and the copula function C , which models the dependence structure on the joint distribution. Drawing from this decomposition, we model the K VPOs observed for a synthetic respondent population by a specific Gaussian copula C^k , $k = 1, \dots, K$.¹⁵ Thus, if respondent i ranks as a member of VPO k , her latent positions X_i^* follow the distribution F^k below:

$$(B.4) \quad F^k(x_1, \dots, x_Q) = C^k(F_1(x_1), \dots, F_Q(x_Q)),$$

where $F_1(x_1), \dots, F_Q(x_Q)$ are the marginal distributions across all respondents for each question q . These marginal distributions are shared by all respondents independently of the construal they belong to. They provide information on the relative frequency of answers of agreement and disagreement.

We "normalize" respondents' positions, translating it to the unit cube. Define $u_{iq}^* = F_q(x_{iq}^*)$ and $u_i^* = (u_{i1}^*, \dots, u_{iQ}^*)$. If respondent i belongs to VPO k , then u_i^* has distribution C^k . That is, u_{iq}^* is uniformly distributed for each question q . An example will help clarify what u_{iq}^* means. Say that the options in question q are ordered from $a_1 = \text{Absolute disagreement}$ to $a_{H_q} =$

15. A Gaussian copula with correlation matrix Σ is defined by

$$(B.3) \quad C(u_1, \dots, u_Q; \Sigma) = \Phi_{\Sigma}(\Phi^{-1}(u_1), \dots, \Phi^{-1}(u_Q)),$$

where Φ is the cumulative distribution function (cdf) of a standard normal variable and Φ_{Σ} be the cdf of a multivariate normal vector with mean zero and covariance matrix Σ . Fix $\Sigma_{qq} = 1$, $q = 1, \dots, Q$, so that the marginal distributions have unit variance. The family of Gaussian copulas is indexed by the correlation matrices Σ . They generalize the Gaussian-like dependence structure to allow for arbitrary (i.e., non-Gaussian) marginal distributions.

Absolute agreement. Then, since u_{iq}^* is uniformly distributed, $u_{iq}^* = 0.7$ means that 70% of the population agrees with the statement of question q less than respondent i . This is what we mean by accurate position of the respondent. The fact that we normalize $u_{iq}^* = F_q(x_{iq}^*)$ makes the position of respondents easier to interpret.

Modeling answer choices. We can also “normalize” the thresholds, so we directly work with u_i^* when generating the data. Let $\bar{\tau}_{iq\eta} = F_q(\tau_{iq\eta})$. Then, according to equation (B.1), respondent i chooses x_{iq} with the following rule:

$$(B.5) \quad x_{iq} = a_\eta \iff u_{iq}^* \in (\bar{\tau}_{iq,\eta-1}, \bar{\tau}_{iq\eta}) \text{ for } \eta = 1, \dots, H_q.$$

The position u_i^* comes from the distribution C^k and is, therefore, subject to the opinion constraints underling VPOs. The random threshold $(\bar{\tau}_{iq\eta})_{\eta=1}^{H_q}$, in turn, model individual variation (idiosyncratic shocks).

The threshold generation we implement takes the neutrality element as its reference point. It also sets up thresholds following symmetry considerations. The threshold set up we establish specifies individual thresholds as position points located increasingly farther away from the neutral element. For instance, consider that the number of answers H_q is odd and there is an explicitly neutrality point. The modeling procedure then draws $\bar{H}_q = (H_q - 1)/2$ ”distances” from this point: $(\omega_{iq\eta})_{\eta=1}^{\bar{H}_q}$.¹⁶

The first distance, ω_{iq1} , determines the length of the ”neutrality interval”, or the range of positions for which a respondent would choose a neutrality element. More formally, we set the interval that determines whether $x_{iq} = a_{\bar{H}_q} = n_q$ (neutrality) to

$$(B.6) \quad (\bar{\tau}_{iq\bar{H}_q}, \bar{\tau}_{iq,\bar{H}_q+1}) = (\bar{n}_q - \omega_{iq1}, \bar{n}_q + \omega_{iq1}).$$

The subsequent distances $(\omega_{iq\eta})_{\eta=2}^{\bar{H}_q}$ measure the length of the rest of the position intervals. We can recursively define the thresholds that diverge farther to the right of the neutrality point by

$$(B.7) \quad \bar{\tau}_{iq,\bar{H}_q+\eta} = \bar{\tau}_{iq,\bar{H}_q+\eta-1} + \omega_{iq\eta} \text{ for } \eta = 2, \dots, \bar{H}_q.$$

16. Distances $(\omega_{iq\eta})_{\eta=1}^{\bar{H}_q}$ can be drawn from arbitrary distributions, as long as one guarantees that $0 < \bar{\tau}_{iq1} < \bar{\tau}_{iq2} < \dots < \bar{\tau}_{iq,\bar{H}_q-1} < 1$.

Then, by symmetry with respect to the neutral position, the thresholds that diverge farther to the left of the neutrality point are:

$$(B.8) \quad \bar{\tau}_{iq,1+\bar{H}_q-\eta} = \bar{\tau}_{iq,2+\bar{H}_q-\eta} - \omega_{iq\eta} \text{ for } \eta = 2, \dots, \bar{H}_q.$$

General algorithm

We specify the algorithm we follow to produce synthetic data based in the discussion above. The parameters that feed the algorithm are the following:

- Set of possible VPOs: $\mathcal{K}$;
- Number of questions: Q ;
- Set of possible answer options for questions: $\mathcal{H}$;
- Interval for neutral position: $\mathcal{N}$.

The algorithm for the production of synthetic databases proceeds as follows. First, we generate the matrix with **latent accurate positions**:

- (1) Draw the number of construals K from $\mathcal{K}$.
- (2) Set the number of questions to Q .
- (3) For each construal k , build the construal correlation matrix Σ^k (see below). The copula C^k defining construal k is the Gaussian copula with correlation matrix Σ^k .
- (4) For each construal k , randomly generate the number of respondents adhering to it: $N_k \in \{100, 101, \dots, 499, 500\}$. All cases are equally probable. The total number of respondents is then $N = \sum_{k=1}^K N_k$.
- (5) For each respondent i ranking as member of VPO k , generate the latent accurate position u_i^* following distribution C^k . Precisely, we generate a random draw $(\xi_{iq})_{q=1}^Q$ from a multivariate normal with mean zero and covariance matrix Σ^k . Then, we set $u_{iq}^* = \Phi^{-1}(\xi_{iq})$ for $q = 1, \dots, Q$.

Then, we generate the **synthetic answer data**:

- (5) For each question q , draw the number of options H_q from $\mathcal{H}$. The answer set is

$$(B.9) \quad \mathcal{A}_q = \{-(H_q - 1)/2, \dots, -1, 0, 1, \dots, (H_q - 1)/2\}.$$

- (6) For each question q , draw neutral position $\bar{n}_q$ from $\mathcal{N}$.
- (7) If thresholds vary at individual level, do nothing.

- (8) For each respondent i and question q :
- (a) To generate thresholds, we compute $\bar{H}_q = (H_q - 1)/2$ distances: $(\omega_{iq\eta})_{\eta=1}^{\bar{H}_q}$. How these are obtained varies depending on whether they are set to be random (see below). Once we have the distances, we get the thresholds using equations (B.6), (B.7), and (B.8).
 - (b) We set answer x_{iq} according to equation (B.5).

APPENDIX C. CORRELATION DISSIMILARITY

Our simulations model issue constraints for each VPO k as a Gaussian copula with a correlation matrix Σ^k . Knowing the specific form of these correlation matrices in our simulated data allows us to ask how good each CEM is at generating estimated VPOS with correlation structures approximating them. An adequate assessment of this questions requires tackling two potential problems: first, that the number of VPOs estimated by the method may differ from the true number of VPOs, and second, that the labelling of each VPOs is arbitrary (i.e., what a certain method labels as “the first construal” may differ from the group that was labelled as “the first construal” when generating the data).

We propose a method to measure how well each method estimates the underlying correlation matrices, which overcomes these issues. The inputs to the method are a collection of “true” and estimated correlation matrices, $(\Sigma^k)_{k=1}^K$ and $(\hat{\Sigma}^k)_{k=1}^{\hat{K}}$, respectively. The number of “true” (K) and estimated ($\hat{K}$) construals may differ. Additionally, the researcher must provide a measure of dissimilarity between two matrices, which we denote by d . We base our assessment of the methods on the Frobenius distance

$$(C.1) \quad d(A, B) = \sqrt{\text{trace}(A - B)(A - B)'},$$

where A and B are matrices and A' denotes the transpose of A .¹⁷ For more information about the Frobenius distance, see (Cohen 2022).

As we have discussed, our measure of performance must deal with the fact that there is no automatic way to pair an estimated VPO with a “true” one. That is, for a given estimated VPKO $\hat{k} \in \{1, \dots, \hat{K}\}$ we do not know the corresponding true VPO $k \in \{1, \dots, K\}$. Thus, we cannot directly compute $d(\Sigma^k, \hat{\Sigma}^{\hat{k}})$. The problem is exacerbated when the number of estimated VPOs differs from the number of true VPOs. To deal with these inconveniences, we will work with all

17. The choice of the Frobenius distance is not particularly relevant. Indeed, in case the dissimilarity measure comes from a norm, the ranking of methods will not depend on the chosen norm. This comes from all norms being equivalent.

possible pairings between $\{1, \dots, \hat{K}\}$ and $\{1, \dots, K\}$. Our measure of performance will then be the minimum of matrix dissimilarity across possible pairings.

We define a *paring* between $\{1, \dots, \hat{K}\}$ and $\{1, \dots, K\}$ as an injective function $p: \{1, \dots, \hat{K}\} \rightarrow \{1, \dots, K\}$. This function assigns a unique “true” VPO $k = p(\hat{k})$ to each estimated VPO $\hat{k}$. Note, however, that no injective function exists when $\hat{K} > K$, that is, when the method overestimates the number of VPOs. To deal with this, we propose to augment the “true” number of VPOs by introducing $\hat{K} - K$ artificial VPOs. The correlation structure of these artificial VPOs is $\Sigma^k = I_Q$, the Q -dimensional identity matrix. That is, issues for these artificial VPOs are not constrained, indicating that the construals are not really present. So, when $\hat{K} > K$, we propose to compute the dissimilarity between the collection of estimated correlation matrices $(\hat{\Sigma}^k)_{k=1}^{\hat{K}}$ and the augmented collection of $\hat{K}$ “true” correlation matrices $(\Sigma^1, \dots, \Sigma^K, I_Q, \dots, I_Q)$.¹⁸

To introduce the measure of dissimilarity between $(\hat{\Sigma}^k)_{k=1}^{\hat{K}}$ and $(\Sigma^k)_{k=1}^K$, consider first that $\hat{K} \leq K$. Denote by P the set of all possible pairings between $\{1, \dots, \hat{K}\}$ and $\{1, \dots, K\}$. Given a dissimilarity measure between pairs of matrices d , the dissimilarity between collections, denoted $\mathcal{D}$, is

$$(C.2) \quad \mathcal{D} \left((\hat{\Sigma}^k)_{k=1}^{\hat{K}}, (\Sigma^k)_{k=1}^K \right) := \min_{p \in P} \left\{ \max_{\hat{k}=1, \dots, \hat{K}} d \left(\hat{\Sigma}^{\hat{k}}, \Sigma^{p(\hat{k})} \right) \right\}, \text{ when } \hat{K} \leq K.$$

As we have discussed, when $\hat{K} > K$, we augment the collection of “true” VPOs by adding VPOs whose correlation structure is the identity. That is, we define the collection $(\tilde{\Sigma}^k)_{k=1}^{\hat{K}}$, with $\tilde{\Sigma}^k = \Sigma^k$, for $k \leq K$, and $\tilde{\Sigma}^k = I_Q$, for $k = K+1, \dots, \hat{K}$. In this case, the dissimilarity between the collections is given by

$$(C.3) \quad \mathcal{D} \left((\hat{\Sigma}^k)_{k=1}^{\hat{K}}, (\Sigma^k)_{k=1}^K \right) := \mathcal{D} \left((\hat{\Sigma}^k)_{k=1}^{\hat{K}}, (\tilde{\Sigma}^k)_{k=1}^{\hat{K}} \right), \text{ when } \hat{K} > K.$$

Remark C.1. An alternative to augmenting the number of “true” VPOs when $\hat{K} > K$ is to consider, in that case, pairings as injective functions $p: \{1, \dots, K\} \rightarrow \{1, \dots, \hat{K}\}$. This leaves, however, some estimated construals unpaired with “true” VPOs. Thus, we believe that this other approach unfairly benefits methods that overestimate the number of VPOs. Indeed, a method can trick this alternative measure by estimating a lot of VPOs, even if some of those are very dissimilar

18. An alternative to this approach would be to consider injective functions $p: \{1, \dots, K\} \rightarrow \{1, \dots, \hat{K}\}$ when $\hat{K} > K$. Below we discuss why we believe that this other approach could be misleading in Remark C.1.

to the “true” VPOs. This is so since the minimum in equation (C.2) will be achieved when the dissimilar estimated VPOs are left unpaired.

A Benchmark to asses method performance

One of the problems of the above measure of dissimilarity is that there is no clear notion of whether a method estimated properly the underlying correlation structure. Certainly, when $\mathcal{D}$ is zero, the estimation is perfect and larger values of $\mathcal{D}$ indicate a worse fit. However, there is no guidance on how to read the magnitude of $\mathcal{D}$. To address this we propose to compare it to a benchmark.

Let $(X_i)_{i=1}^N$ be a synthetic dataset generated as described in Appendix B. For each synthetic observation i , we know the VPO it belongs to. That is, if there are K “true construals” in the synthetic data, we observe the “true” VPO membership function $c: \{1, \dots, N\} \rightarrow \{1, \dots, K\}$, where $c(i) = k$ means that observation i belongs to construal k . With this information, we can estimate $\bar{\Sigma}^k$ using only the observations i such that $c(i) = k$. This correlation matrix differs from the “true” correlation matrix for construal k (Σ^k) for two reasons: (i) there is some sample variation and, more importantly, (ii) we only observe answers X_i and not the accurate latent positions X_i^* .

When we estimate the correlation matrices $(\hat{\Sigma}_k)_{k=1}^{\hat{K}}$ using a given method, we also face the above two sources of noise. However, on top of them, each method must estimate the construal membership function. Therefore, we believe that $\mathcal{D}((\bar{\Sigma}^k)_{k=1}^K, (\Sigma^k)_{k=1}^K) := \mathcal{D}^b$ provides a good benchmark to measure the performance of the method when estimating the underlying correlation structure.

Once we compute $\mathcal{D}^b$, we report

$$(C.4) \quad \frac{\mathcal{D}\left((\hat{\Sigma}^k)_{k=1}^{\hat{K}}, (\Sigma^k)_{k=1}^K\right)}{\mathcal{D}^b},$$

where $(\hat{\Sigma}_k)_{k=1}^{\hat{K}}$ are derived using each method. The magnitude of the above ratio is meaningful. For instance, if the ratio is 2 for a given method, one can say that the “estimation of the underlying correlation structure with the method is twice as bad as if we knew the true membership of each observation”.