

Online Appendix

The Ideologies of Organized Interests & Amicus Curiae Briefs

Large-Scale, Social Network Estimation of Ideal Points

Sahar Abi-Hassan* Janet M. Box-Steffensmeier[†] Dino P. Christenson[‡]
Aaron R. Kaufman[§] Brian Libgober[¶]

Contents

1	Mathematical Notation and Estimation Procedure	1
2	Egocentric Networks Show Mixed Strategies	3
3	Interest Group Consistency Before the Court	5
4	Comparison to Network Cohesion Model	5
5	Comparison to Latent Network Model	7
6	Bootstrapped Standard Errors	8
7	Benchmarking Validity Tests	9
7.1	Comparing When Measures Disagree	10
7.2	Network Simulation Results	10
8	Comparison of Training Labels	13
9	Drawbacks and Limitations	14
10	A Data-Generating Process	15
10.1	The Game	15
10.2	The Assumptions	16

*Assistant Professor, Department of Public Policy and Political Science, Mills College, sabihassan@mills.edu

[†]Vernal Riffe Professor of Political Science and Sociology, The Ohio State University, box-steffensmeier.1@osu.edu

[‡]Corresponding author. Associate Professor, Department of Political Science, Boston University, dinopc@bu.edu

[§]Assistant Professor, Division of Social Sciences, New York University, Abu Dhabi aaronkaufman@nyu.edu

[¶]Postdoctoral Associate and Lecturer, Yale University, brian.libgober@yale.edu

1 Mathematical Notation and Estimation Procedure

Let there be a set of amicus-writing organizations A that have written a set of briefs B , some co-authored and some not. We think of A and B as together being a set of nodes constituting a bipartite graph G . Let $p(v, 0)$ be the initial ideal point of each vertex in the graph. Let $N(v)$ be the neighborhood of v . Then we assume that the co-signers of brief b (formally, $N(b)$) negotiate the content of the brief so that its position reflects the average of their preferred positions, perhaps with some noise $p(b, 1) = \frac{1}{|N(b)|} \sum_{v \in N(b)} p(v, 0) + \epsilon_{b1}$. Organizations come in two types. Some organizations have “known” positions, in which case $p(a, 1) = p(a, 0) = p(a)$ is fixed and constant. Alternatively, some organizations a have “unknown” positions. In this case, we assume that the average position of the briefs a signed is a noisy estimate of its position. More formally, $p(a, 1) = \frac{1}{|N(a)|} \sum_{v \in N(a)} p(v, 1) + \epsilon_{a1}$ where $\epsilon_{a1} \sim N(0, \sigma_1^2)$. Since $p(a, 1)$ is a noisy estimate of a ’s position, $p(b, 1)$ no longer reflects the average position of the organizations that co-signed it. To address this issue, we focus on finding the expected “steady state” of this system, where the updating from $t - 1 \rightarrow t$ follows the same rule as from $0 \rightarrow 1$. In such a state, all briefs reflect a fair bargain between co-signers based on their ideology and the average of the briefs an organization signs is a noiseless estimate of their ideology. In fitting the estimates, we may ignore the noise which will eventually drop out, and prefer to calibrate the uncertainty associated with these estimates through bootstrapping and cross-validation.

To calculate the state, note first that we can formalize the brief update step through a matrix $\mathbf{B}$ whose entry b_{ij} is the number of briefs co-signed by i and j divided by the total number of briefs that i has signed, unless $i = j$ in which case $b_{ij} = 0$. Let $\vec{p}$ be a vector where p_i is i ’s original ideology score if available.¹ Then the value after the first brief step is $\mathbf{B}\vec{p} + \vec{\epsilon}_{1b}$. Similarly, we can construct a matrix $\mathbf{O}$ to describe the organizational update step, and note that the estimate of ideology after the org step is $\mathbf{A}\mathbf{B}\vec{p} + \mathbf{A}\vec{\epsilon}_{1b} + \vec{\epsilon}_{1o}$. The scores evolve as follows:

$$\begin{aligned} & \mathbf{A}\mathbf{B}\mathbf{A}\vec{p} + \mathbf{A}\mathbf{B}\mathbf{A}\vec{\epsilon}_{1b} + \mathbf{A}\mathbf{B}\vec{\epsilon}_{1a} + \mathbf{A}\vec{\epsilon}_{2b} + \vec{\epsilon}_{2a} \\ & \mathbf{A}\mathbf{B}\mathbf{A}\mathbf{B}\vec{p} + \mathbf{A}\mathbf{B}\mathbf{A}\mathbf{B}\vec{\epsilon}_{1b} + \mathbf{A}\mathbf{B}\mathbf{A}\mathbf{B}\vec{\epsilon}_{1a} + \mathbf{A}\mathbf{B}\mathbf{A}\vec{\epsilon}_{2b} + \mathbf{A}\mathbf{B}\vec{\epsilon}_{2a} + \mathbf{A}\vec{\epsilon}_{3b} + \vec{\epsilon}_{3a} \\ & (\mathbf{A}\mathbf{B})^n \vec{p} + \sum_{i=1}^n \{ (\mathbf{A}\mathbf{B})^{n-i} \mathbf{A}\vec{\epsilon}_{ib} + (\mathbf{A}\mathbf{B})^{n-i} \vec{\epsilon}_{ia} \} \end{aligned}$$

Under suitable assumptions about the variance of the noise as the system evolves, $\lim_{n \rightarrow \infty} \sum_{i=1}^n \{ (\mathbf{A}\mathbf{B})^{n-i} \mathbf{A}\vec{\epsilon}_{ib} + (\mathbf{A}\mathbf{B})^{n-i} \vec{\epsilon}_{ia} \}$ is also Gaussian with finite variance and expected value 0. Therefore, we may focus our attention on the value of $\lim_{n \rightarrow \infty} (\mathbf{A}\mathbf{B})^n \vec{p}$, which is the expected steady state of the system.

To see how we may calculate this value analytically, let $\mathbf{M} = \mathbf{A}\mathbf{B}$ and note that $\mathbf{M}$ is a valid transition matrix for a Markov process on the digraph of briefs and organizations. Nodes with origin scores are “absorbing states,” in the sense that if one were to transition to these states one

1. If i has no origin score available, p_i can be anything

would never leave. All other nodes represent transient states. If we reorder the nodes so that organizations with origin scores are first, we may write

$$\mathbf{M} = \begin{bmatrix} \mathbf{I} & 0 \\ \mathbf{S} & \mathbf{Q} \end{bmatrix}$$

Here $\mathbf{S}$ contains the transition probabilities of an node without a score to a node with an origin score, while $\mathbf{Q}$ contains the transition probabilities between nodes lacking scores. Put differently, $\mathbf{S}$ represents the probabilities of going to an absorbing state on the digraph when starting from a transient state, while $\mathbf{Q}$ represents the transition probabilities of going from one transient state to another. Further, block matrix multiplication shows

$$\mathbf{M}^k = \begin{bmatrix} \mathbf{I} & 0 \\ \mathbf{S}(\mathbf{I} + \mathbf{Q} + \mathbf{Q}^2 + \dots + \mathbf{Q}^k) & \mathbf{Q}^k \end{bmatrix}$$

Let us describe what happens to each block in turn as $k \rightarrow \infty$. Clearly, the probability of being outside an absorbing state goes down with k , and must eventually become 0.² As a result, $\lim_{k \rightarrow \infty} \mathbf{Q}^k = 0$. Further we have the identity $(\mathbf{I} + \mathbf{Q} + \mathbf{Q}^2 + \dots + \mathbf{Q}^k + \dots)(\mathbf{I} - \mathbf{Q}) = \mathbf{I}$ therefore $\mathbf{S}(\mathbf{I} + \mathbf{Q} + \mathbf{Q}^2 + \dots + \mathbf{Q}^k + \dots) = \mathbf{S}(\mathbf{I} - \mathbf{Q})^{-1}$, if $(\mathbf{I} - \mathbf{Q})$ is invertible. Since the eigenvalues of $\mathbf{Q}$ are strictly less than 1 for $\lim_{k \rightarrow \infty} \mathbf{Q}^k = 0$, the inverse of $(\mathbf{I} - \mathbf{Q})$ must exist. Therefore we have

$$\mathbf{M}^\infty = \begin{bmatrix} \mathbf{I} & 0 \\ \mathbf{S}(\mathbf{I} - \mathbf{Q})^{-1} & \mathbf{0} \end{bmatrix}$$

Let $\tilde{p} = \mathbf{M}^\infty \tilde{\mathbf{p}}$. Then $\mathbf{M}\tilde{p} = \mathbf{M} \cdot \mathbf{M}^\infty \tilde{p} = \mathbf{M}^\infty \tilde{p} = \tilde{p}$ so the weighted average of neighbor property will obtain for those nodes which are not given exogenously. Further, by construction, $\tilde{p}_i = \tilde{p}_i$ if i is an absorbing state/origin node.

Clearly, the term $\mathbf{S}(\mathbf{I} - \mathbf{Q})^{-1}$ is the key one for estimation. It reflects the weights ultimately given to each of the nodes with origin scores for each of the nodes we estimate. Potentially, these weights are also of interest, and could be useful to researchers. For example, if the exogenous ideal point measures are not estimated with certainty, then they can be used to efficiently impute confidence intervals around the endogenous measures, for example by sampling. Alternatively, they could be used for sensitivity analysis. Another important point about the formulation above is that we have implicitly assumed that all nodes without origin scores were connected by some degree to one with an origin score. If it were not, then we would have no basis for ever imputing its ideology purely on the basis of its neighbors. Finally, the block nature of the matrix $\mathbf{M}^\infty$ gives some insight into why it does not matter what values are initially supplied for the nodes without origin scores. In the expected steady state, the estimated scores depend only on the origin scores.

2. We assume that all transient nodes are connected somehow to an absorbing node. A node that is not connected to some node with an origin score is one to which we are unable to assign an ideology score, and so we may eliminate all these prior to estimation.

2 Egocentric Networks Show Mixed Strategies

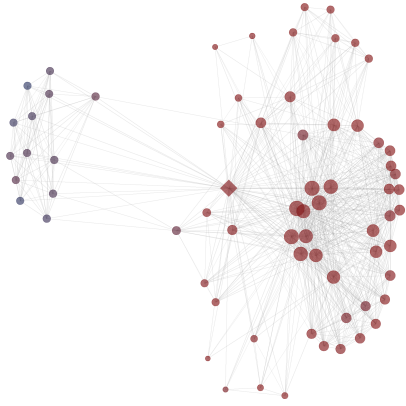
We can look more closely into the coalition structure of particular interest groups with their egocentric networks, or egonets. Egonets present all the organizations that have cosigned with an organization, the ego, and the cosigning links between those organizations. Thus, they can provide insights into the diversity of strategies particular nodes play within their own networks. While interest groups have been shown to employ mixed strategies before the Court—for example, acting as equal teammates, or leaders, as well as avoiding coalitions altogether (Box-Steffensmeier and Christenson 2014; Box-Steffensmeier et al. 2018)—the role of ideology in these coalitions are largely unexplored.

In Figure 1, we observe six prominent amicus cosigning organizations: two conservative (Gun Owners of America and the Free Speech Defense and Education Fund), two moderate (New York Times and Johnson & Johnson Co.), and two liberal (Feminist Majority Foundation and NARAL Pro-Choice America). As in all the figures, all of these organizations’ ACNet scores were imputed based on the network. We also provide these organizations’ ideal points and the average ideal points of their ego networks in Table 1. The first column indicates the ideology or ACNet score of the ego. The second column indicates the average ideology of all organizations in the organization’s ego network. Substantial differences between the ego’s ideology and their cosigners’ ideologies suggest an ideologically heterogeneous strategy, whereas similar scores suggest a homogenous one. The final row indicates the average ideology of all organizations in the data, and the *degree-weighted* average ideology of all organizations in the data, to give an idea of how different these six groups are from the rest of the network.

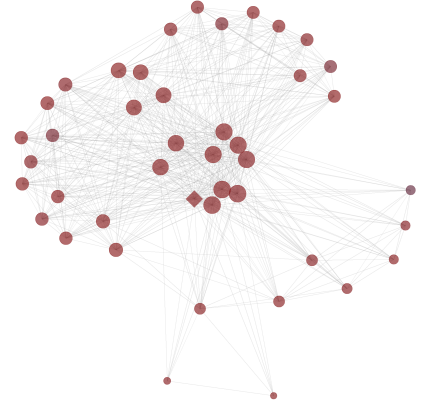
	Organization Ideology	Ego Network Ideology
Gun Owners of America	0.67	0.85
Free Speech Defense and Education Fund	0.89	0.95
New York Times	0.15	0.09
Johnson & Johnson	0.13	0.16
Feminist Majority Foundation	-0.61	-0.67
NARAL Pro-Choice America	-0.62	-0.60
Overall Average	-0.18	-0.30

Table 1: Ideologies of Ego-Networks of Prominent Amicus-Filing Organizations

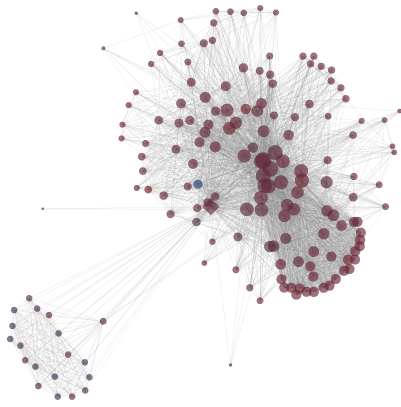
Examination of these six ego networks reveals important facets of their organizations’ legal strategies. Foremost—and just as it appeared in the full network plot (Figure 1)—the egonets show that interest groups generally work with organizations of similar ideological backgrounds. The ego ideologies in Table 1 are similar to the average in their networks. Of course, Figure 1 also shows that some of these groups work in more ideologically heterogeneous networks than others. The Feminist Majority Foundation and the Free Speech Defense and Education Fund primarily cosign with dense, highly connected groups of co-ideological organizations. These groups largely work as equal teammates, situated among groups that work with one another and agree



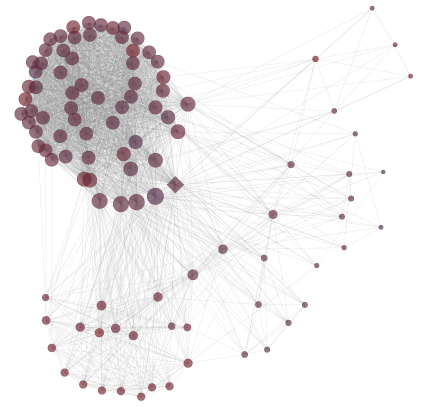
(a) Gun Owners of America



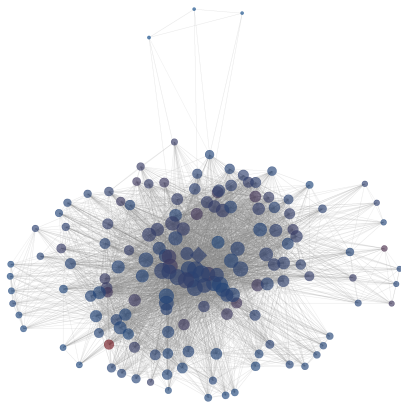
(b) Free Speech Defense and Education Fund



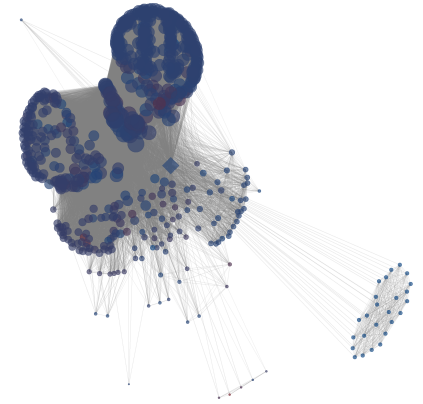
(c) New York Times



(d) Johnson & Johnson Co.



(e) Feminist Majority Foundation



(f) NARAL Pro-Choice America

Figure 1: Ego-Networks of Prominent Amicus-Filing Organizations. *Circles indicate organizations. Diamonds indicate the ego node. Nodes are colored by ACNet score from liberal (blue) to conservative (red). Nodes are sized by their degree centrality.*

ideologically. By contrast, the Gun Owners of America and NARAL Pro-Choice America have two almost wholly disconnected subgroups, meaning that they act as leaders and information brokers, bringing together groups that would be otherwise disconnected. The Gun Owners of America has a subnetwork of conservative cosigners and one subnetwork of strongly liberal co-signatories. Of the two clusters of NARAL Pro-Choice America cosigners, one is a diverse array of liberal interest groups while the other consists exclusively of state-specific NARAL chapters. Located at the center of the ideological spectrum, the *New York Times* and Johnson & Johnson Co both cosign with a diverse and interconnected group of moderate organizations from across the political spectrum. Ultimately, the egonets not only suggest support for mixed coalition strategies, but also that few exist in perfectly homogenous ideological networks. While polarization is the norm, moderate groups seek out partnerships with ideologically diverse groups from time to time, creating a relatively dense network of the population of interest groups.

3 Interest Group Consistency Before the Court

How consistently ideological are interest groups in their amicus curiae signing behavior? We combine the Supreme Court Database’s manual coding of court decision ideology and the Amicus Curiae Networks Project’s data set of organizations’ amicus signing, including whether the brief was in support of the petitioner or the respondent, to study this question.

For each organization, we tabulate the number of times they signed or cosigned a brief advocating for the liberal position and the conservative position. Next, we calculate each organization’s modal ideological position, and the proportion of briefs they sign of the *opposite* position. We find that across all organizations, about 15% of signings are not in support of an organization’s usual ideological position. For example, the ACLU has signed 530 briefs for the liberal side, but 82 for the conservative side (14.4%). The National Council of La Raza has signed 26 liberal briefs and 2 conservative ones (7%); the National Sheriffs Association has signed 33 conservative briefs and 2 liberal ones (6%).

Some of the more interesting organizations are media organizations: CBS Inc., the Hearst Corporation, and the Recording Industry Association of America have all signed as many liberal as conservative briefs (7, 9, and 6 briefs each, respectively). This is likely related to the fact that First Amendment cases are often ideologically ambiguous. We intend to pursue focused explorations of this nature in subsequent work.

4 Comparison to Network Cohesion Model

As an alternative and complimentary approach to our iterative weighted average approach, we also consider a network cohesion model (Li, Levina, and Zhu 2016). Generally, this model discourages network-proximate nodes for having dissimilar outcomes; it provides a general framework for incorporating network linkage into predictive models alongside more standard covariates, and

can substantially improve predictive performance in the presence of strong network cohesion. We measure the internal validity of this model through edge-sampled cross-validation.

The intuition for this procedure is that if we assume network homophily and we know *a priori* node characteristics, we can predict which nodes might have edges between them. Therefore if we randomly remove some number of edges, a well-fitted model should reliably be able to predict where those edges were.

The key tuning parameter λ in the network cohesion model controls the weight put on network cohesion, where higher values constrain the model to more highly prefer homophily. We select the optimal λ through our edge-sampled cross-validation procedure. In the absence of homophily, we would expect that the cross-validation procedure would yield substantially similar error rates for all values of λ ; we find that the optimal value of λ is 0.9, lending support to our homophily assumption.³

While the network cohesion model can flexibly incorporate covariates, we experiment with doing so and choose not to for both computational and theoretical reasons. We find that including covariates in the imputation model improves cross-validation accuracy only at the margin, reducing error by about 0.5%, from an MSE of 0.599 to an MSE of 0.596, and at the cost of computational efficiency.

Secondly, incorporating some of those covariates could subsequently create endogeneity concerns should we choose to explore the relationship between our network-based ideology measure and other case-level outcomes. Our imputed data are underdispersed relative to the training labels: the interquartile range is 0.34 in our imputations compared to 1.35 in the training data. We directly estimate this underdispersion through cross-validation, and rescale our predictions by multiplying by the ratio of the standard deviation of predicted ideology to the standard deviations of the true ideology as measured by DIME.

In choosing between the network cohesion model and the iterative weighted average approach to imputing organization ideal points, we consider the methods' internal validity, flexibility, and potential for bias in secondary analysis. For all three reasons we prefer the iterative weighted average model.

Internal validity. Though the network cohesion model includes a complementary cross-validation method, it is possible to cross-validate the iterative weighted average method as well, its predictions are equally stable. Moreover, estimated ideal points from both models correlate in excess of 0.91.

Flexibility. The network cohesion model includes a tuning parameter to adjust for assumptions of network homophily, while the iterative weighted average model requires no tuning parameters.

3. Standard cross-validation builds a model on a subset of the observations for which the outcome is known, then predicts the outcome for a held-out test set; the accordance between the model's predictions and the true outcomes is taken as a measure of predictive accuracy. Network data is less amenable to this procedure, as the observation unit can be nodes, edges, or both. We proceed using edge-sampled cross-validation, but our results are substantially identical using a more conventional node-sampled approach.

Second-stage analysis. New measures, while interesting descriptively, are most useful in second-stage analysis as either the dependent or independent variable in a regression analysis. Here the dispersion problem of the network cohesion model becomes prohibitive: if we cannot reliably compare political elites to interest groups due to a dispersion incongruity, our estimates’ usefulness is severely diminished.

5 Comparison to Latent Network Model

The R library `latentnet` is a popular tool for network embedding and clustering (Krivitsky and Handcock 2008). We believe our semi-supervised approach offers a number of advantages over it, but as it is the most widely used tool for related purposes, in this section we compare its performance to our own model’s on a number of dimensions. First we will discuss what we consider to be our method’s key advantages, then we will offer our benchmark comparisons. Most of these advantages arise from the semi-supervised nature of our approach: our method incorporates a flexible amount of training data for some but not all nodes in the network, including training data in only one of the two components of a bigraph.⁴

Computational efficiency Running the simplest version of a `latentnet` model proved computationally intractable on this data set – a roughly 300Mb network of 15,000 nodes; computation crashed after two weeks on a laptop with 32Gb of RAM and an Intel i-9 16-core 2.4Ghz processor. A version with 4,000 nodes and 15Mb took more than 5 days to complete. The full version of our iterative weighted averages approach takes less than a second.

Estimation quality Our ACNet scores correlate with the posterior mode estimates of a `latentnet` model at 0.14, suggesting that both network-based tools pick up some in-common signal. While there is no gold-standard scores to which we can compare both sets of estimates, there are empirical regularities we can examine to validate both measures. In Figure 2 we show that while our ACNet scores tightly cluster the various state and local branches of the American Civil Liberties Union on the liberal side of the spectrum, the `latentnet` scores place some ACLU branches on both extremes of the ideological spectrum.

Common scaling Our model imputes ideal points on the same scale as its training labels, which are Bonica’s DIME scores. As a result, our scores are on the same scale. Since it is not possible to include training labels in `latentnet` models based on our communications with the library’s authors, we cannot enforce as such. Our scores range, roughly, from -3 to +3; `latentnet` scores range from -5,000 to +11,000, which is not an interpretable scale and is not comparable to other well-studied actors.

4. In correspondence with the authors of the `latentnet` library, we were informed that this semi-supervised approach is not possible within their library’s framework.

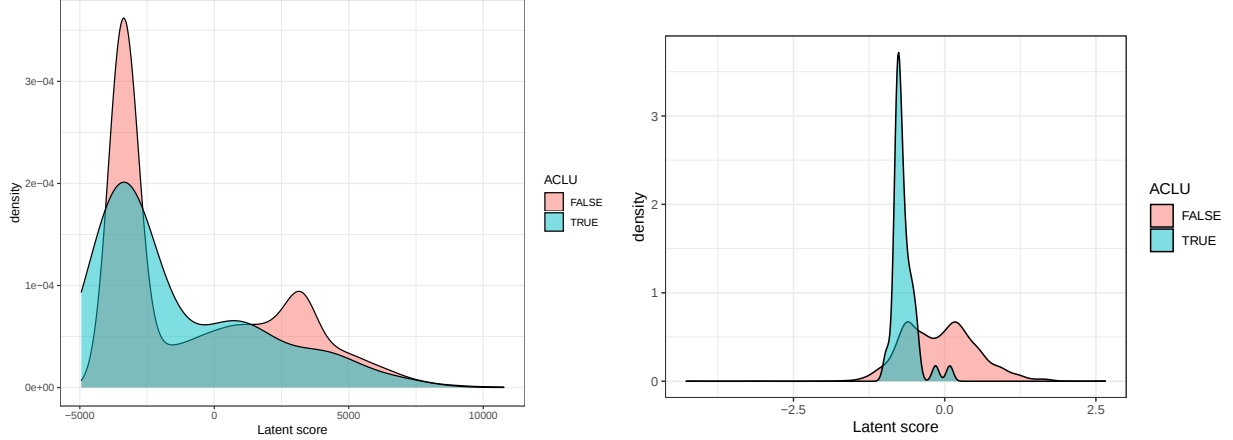


Figure 2: The Latentnet Distribution Scores for ACLU and non-ACLU organizations.

Bigraph estimation The latentnet library does not offer any intrinsic tools for defining bigraph estimation. The ideal matrix for estimation is an $i \times j$ matrix where i indexes organizations and j indexes briefs; latentnet can handle this matrix when it is coerced into an $i + j \times i + j$ matrix, at the loss of computational efficiency. On the other hand, our iterative weighted averages approach leverages this bigraph structure by iteratively updating brief weights then organization weights.

6 Bootstrapped Standard Errors

We calculate bootstrapped standard errors for our ideal points holistically. The source of uncertainty in our scores is not the network, which is fixed, but rather from the source scores we derive from Bonica (2013) and from the estimation procedure itself. Therefore, to produce bootstrapped standard errors, we follow three steps.

First, we calculate uncertainty intervals associated with DIME scores using the data underlying DIME scores themselves. For each organization, we obtain the complete set of contributions they make, the dollar amounts associated with those contributions, and the ideal points of the recipients of those donations. We calculate the uncertainty for a DIME score as the standard error of the weighted average of their donation targets’ ideal points.

Second, to complete a single bootstrap iteration, we draw from the distribution implied by the point estimate and standard error for each organization. That is, if an organization’s DIME score is 1 with a standard error of 0.1, we draw a new value for that organization’s ideal point from a Normal distribution with mean of 1 and standard deviation of 0.1. We do this for each training organization in our sample, and repeat the entire procedure for 2,000 bootstrap iterations. This produces an $n \times p$ matrix where p is 2,000 columns and n is equal to the total number of organizations and briefs in our network.

Finally, for each organization and brief, we calculate the standard deviation of their scores across the set of 2,000 bootstrap iterations.

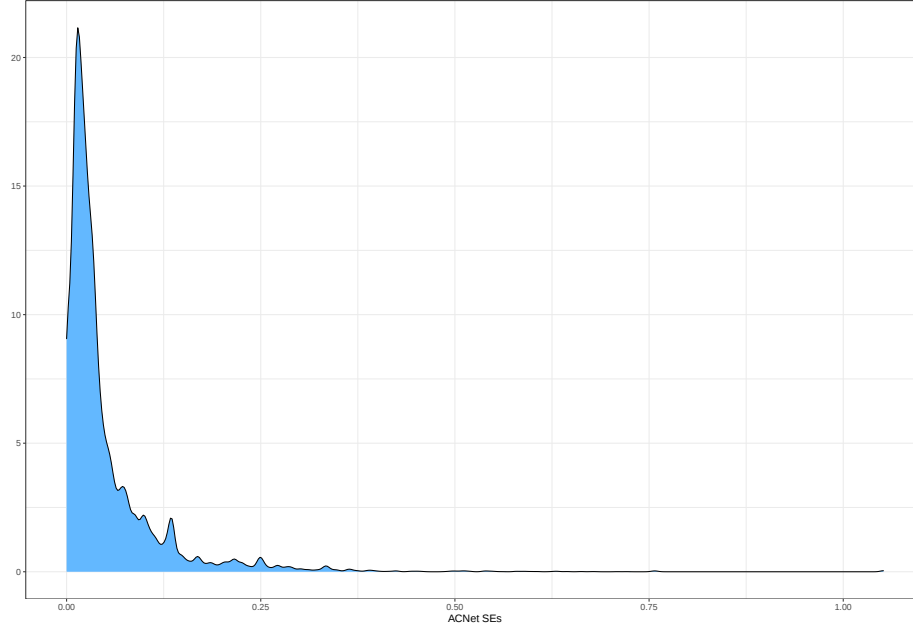


Figure 3: Density plot of bootstrapped standard error estimates.

Our standard errors are generally small, with median of 0.02 and a mean of 0.04, though the max is 1.06. These largest standard errors tend to be organizations with few amicus briefs co-signed, extreme scores, and off-axis interests: Free Speech Advocates, the National Business Aircraft Association, and Georgetown University all have standard errors greater than 0.5. We display the density in Figure 3.

7 Benchmarking Validity Tests

How confident should we be in light of our cross-validation results? We benchmark our results by leveraging the fact that many organizations in the DIME database have multiple entries. We suppress this multiplicity in our estimation procedure by taking a weighted average of scores with weights determined by the number of reports used for the estimation. Nevertheless, to evaluate the reliability of our measures we examine the relationship between the various scores in the DIME database for any single organization. 1,385 organizations have at least two distinct scores within the DIME database, allowing us to calculate similar performance statistics as in Table 4. For example, if we observe two entries for “JP Morgan Chase” in the DIME database, we can calculate both the difference in DIME scores between those entries, and whether those scores are on the same side of zero. Table 2 collects these statistics. The mean absolute deviation between two scores for the same amicus-signing organization in DIME is 0.656, the correlation is 0.373, and the dichotomous classification is the same 64% of the time. That is, our cross-validation procedure predicts DIME scores better than other DIME scores: it produces lower mean absolute deviation, and produces scores on the same side of zero more regularly.

No. Orgs	No. Unique Score Orgs	No. Pairs	Same Name Corr.	MAD	Side Agreement
1,982	597	116,102	0.341	0.681	0.654

Table 2: Examination of Duplicate Scores in DIME

This is less surprising than it sounds: many of the DIME estimates are measured with error, and by leveraging the network we add exogenous information to the signal from DIME. An important source for the disagreement is that the DIME dataset includes record linkage inaccuracies whereby organizations are linked to incorrect FEC contributions. Another reason why the estimates might differ is that interest groups may engage in heterogeneous behavior, in so far as they may behave differently with regards to their political donation activity than they do in amicus cosigning—a possibility we address in more detail below. For both these reasons we neither expect nor desire a perfect accuracy with the DIME scores.

7.1 Comparing When Measures Disagree

Our cross-validation finds cases where network-based and campaign finance-based ideal point estimates diverge substantially. In these cases, which scores have stronger face validity?

In all cases that we examine, ACNet scores appear to have greater face validity than DIME scores. The most extreme differences are for the newspaper USA Today, a moderate newspaper, which we score at 0.18 and DIME scores estimate to be extremely liberal. Other notable disagreements are Clif Bar, which DIME scores as conservative but we score as liberal. Influence Watch describes Clif Bar as left-leaning, primarily advocating for environmental issues and supportive of public healthcare. Hospitality House is a rehabilitation clinic for drug and alcohol abuse that specializes in “evidence-based principles of treatment and recovery, together with multiple treatment phases that slowly reflect the increased levels of personal and social responsibility”; National Voter Outreach Inc is not a ballot access organization, but a petition management company best known for supporting Florida’s 2008 constitutional amendment to lower property taxes and for Republican Sue Lowdon’s candidacy for Lieutenant Governor in Nevada in 2010.

7.2 Network Simulation Results

In the main text and in Table 5 in this Appendix, we present accuracy metrics for our method to estimate ideal points in the amicus curiae network. In another attempt to benchmark these accuracy measures, we conduct a pair of simulation studies designed to bound the predictive accuracy we might expect our method to produce.

In our main analysis, we observe nearly 15,000 nodes of the amicus curiae co-signing network, but we only observe a (noisy) measure of ideology for almost 3,000 of them, approximately one-fifth. We design a cross-validation procedure (Section 2.3) whereby we withhold ideal points for a subset of those 3,000 nodes and then use our tool to predict those scores. In comparing the predicted scores to the truth, we can assess our method’s accuracy.

Organization	ACNet	DIME
USA TODAY	0.18	-4.68
AUTISM NATIONAL COMMITTEE INC	-0.40	-4.73
ASSOCIATION OF COMMUNITY ORGANIZATIONS FOR REFORM NOW	-0.61	-4.21
CITIZENS IN CHARGE FOUNDATION	0.41	3.60
INSTITUTE FOR HUMANIST STUDIES	-0.49	2.65
ENVIRONMENT MAINE	-0.49	-3.51
HOSPITALITY HOUSE INC	-1.88	1.02
POPULATION-ENVIRONMENT BALANCE	-0.65	2.12
CLIF BAR AND CO	-1.70	0.99
TEXANS FOR PUBLIC JUSTICE	-0.60	-3.29
SOUTH DELTA WATER AGENCY	-1.34	1.09
ORGANIZATION OF CHINESE AMERICANS INC	-0.62	1.70
CATO INSTITUTE	0.35	2.65
NATIONAL VOTER OUTREACH INC	1.20	-1.08
PROFESSIONAL SERVICES COUNCIL	1.06	-1.15
GOVERNMENT ACCOUNTABILITY PROJECT	-1.00	1.20
MISSIONARY OBLATES OF MARY IMMACULATE	-0.52	1.66
INDEPENDENCE MINING COMPANY INC	0.82	-1.33
CASCADE POLICY INSTITUTE	0.81	-1.32
RITA JOHNSON	-0.61	1.51
GRANITE BROADCASTING CORPORATION	0.79	-1.25
CENTER FOR RESPONSIBLE LENDING	-0.60	-2.63
ARIZONA WILDLIFE FEDERATION	-0.83	1.16
CAMPAIGN FOR A COLOR-BLIND AMERICA	1.52	-0.48
ARISE RESOURCE CENTER	-0.84	1.15

Table 3: ACNet Scores & DIME Scores

We find that our method has a mean absolute deviation (MAD) of 0.6 and predicts an ideology on the correct side of zero in 74% of the cases. But holding aside the caveats discussed about the quality of underlying DIME data, are those results strong enough to give us confidence that our method is performing well? To assess whether these results are strong, we conduct two simulation studies to obtain comparable measures of MAD and our same-side-of-zero measure. For both of these simulations we observe (by construction) the ideal points of every node in the network.

We perform simulations on two sets of networks: legislative cosponsorship in the 117th Congress⁵, and a fully simulated network. While the legislative network captures some of the real-world patterns we can expect from the amicus network, there are features other than ideology that predict cosponsorship, such as personal networks or belonging to the same state delegation. For these reasons, we design our second network to form edges entirely on the basis of ideological distance. The first simulation gives us a realistic expectation for how strong our results could be in the presence of very strong homophily, while the second simulation uncovers what scores we can expect under weaker homophily but where only ideology plays a role in edge formation.

5. We derive this network loosely from <https://github.com/briatte/congress>.

The legislative cosponsorship network has 447 legislators as nodes, and legislative cosponsorship as edges. We also observe each legislator’s DW-NOMINATE score. We model this network in the `ergm` library where edges are a function of each legislator’s party, a legislator dyad’s party match, and the distance in their DW-NOMINATE scores. Then, we draw from this model to produce hundreds of simulated legislative cosponsorship networks.

The fully simulated network begins by drawing 1,000 briefs and assigning them ideal points from a bi-modal mixture of Gaussians, then drawing 1,000 organizations and assigning them ideal points as well. Then, each brief draws a number of co-signers from an exponential distribution. Finally, each brief samples from the set of 1,000 organizations to fill its co-signers where the probability of selecting organization i is equal to the inverse of the absolute difference between the brief’s ideal point and the organization’s ideal point⁶. This mimics the observed amicus network because there are many organizations with few amici and few organizations with many. We produce hundreds of simulated models according to this data generating process.

Next, we perform our validation exercise. Fully observing each simulated network’s ideal points, we obscure a varying proportion of those ideal points ranging from observing only 5% of the ideal points to observing 95%, nearly the entire network. Then, having forgotten some of the ground truth ideal points, we use our method to estimate those scores and compare our predictions to the hidden ground truth. We present these results in Figure 4, with results from the 117th Congress on the top panel and the fully synthetic network on the bottom.

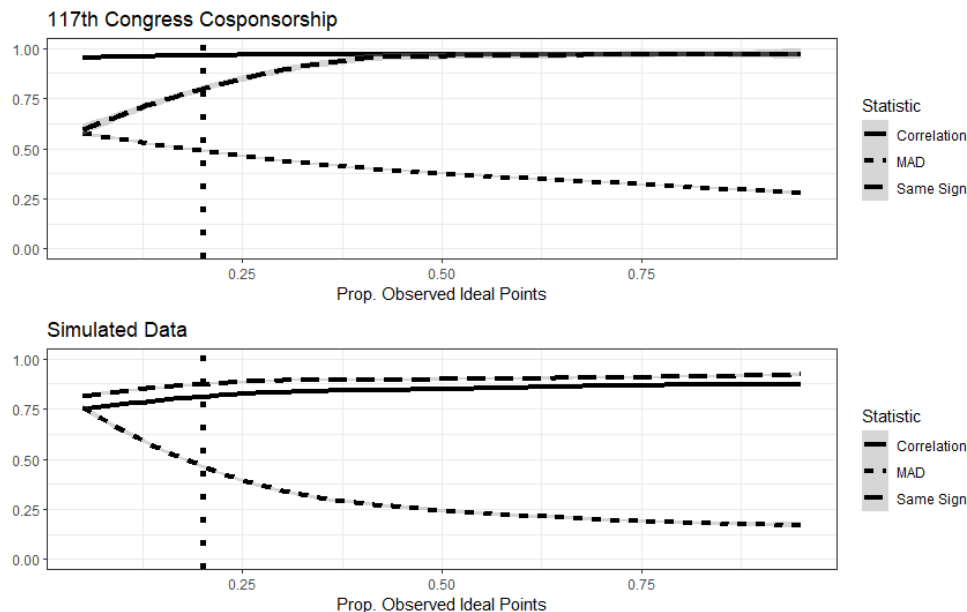


Figure 4: As the proportion of nodes with observed ideal points increases, three measures of accuracy all improve.

We observe that all three measures of accuracy (correlation, MAD, and same-side) improve as

6. We experimented with the squared distance, and the distance to the 1.5 exponent, but these networks produce such extreme homophily that our method performs nearly perfectly every time.

the proportion of observed nodes increases. This makes intuitive sense: as we include more data, the method has more information with which to infer missing nodes’ ideal points. The dotted black line at $x = 0.2$ corresponds to the proportion of observed nodes in our amicus curiae cosigning network. For both panels, the MAD at $x = 0.2$ is approximately 0.5, qualitatively close to our observed MAD of 0.5; as well, the same-side statistic is approximately 0.8 for both, quite similar to our observed 0.74. The two simulations primarily diverge in the correlation statistic (which we do not calculate for the amicus curiae cross-validation). Overall, the new simulations suggest a great deal of confidence in our approach and resulting scores.

8 Comparison of Training Labels

How sensitive are our ideal point estimates to the choice of scores we use to seed the method? To address this, we conduct our estimation procedure replacing the Bonica scores with five alternative sets of scores: Crosson, Furnas, and Lorenz 2020, Hansford, Depaoli, and Canelo 2016, Barberá 2015, and two ensemble measures (Kaufman, King, and Komisarchik 2021): a rowwise average of the above three measures, omitting missing values; and rowwise average of above three measures and DIME scores, omitting missing values⁷.

Crosson, Furnas, and Lorenz 2020 use interest groups’ public position-taking on bills to score approximately 2,600 organizations; Hansford, Depaoli, and Canelo 2016 use amicus curiae briefs to score about 600 organizations; and Barberá 2015 estimates the scores of Twitter users based on their “linking” behavior (see Table 4). Among common observations, the three measures correlate with each other at an average of 0.58; and their correlations to DIME scores average 0.53 as show in Figure 5.

Citation	Training Orgs	Scorable	Designed for Organizations
Bonica (2013)	2,992	24,280	Yes
Crosson et al. (2020)	1,037	22,130	Yes
Hansford & Depaoli (2019)	548	22,203	Yes
Barberá (2015)	481	22,141	No

Table 4: Measures of Interest Group Ideology

To ensure an apples-to-apples comparison, for each measure we find the organizations scored by that measure that are also scored by Bonica (2013). Next, we perform leave-one-out cross-validation using that measure’s scores for only those in-common organizations, and perform leave-one-out cross-validation using DIME scores for only those in-common organizations. Then we compare cross-validation results for that pair. We present these results in Table 5. Since there are fewer labeled organizations for these three methods, we note that using them generates far less scorable organizations than when using DIME scores.

7. These last two methods are ensemble methods as per Kaufman, King, and Komisarchik 2021.

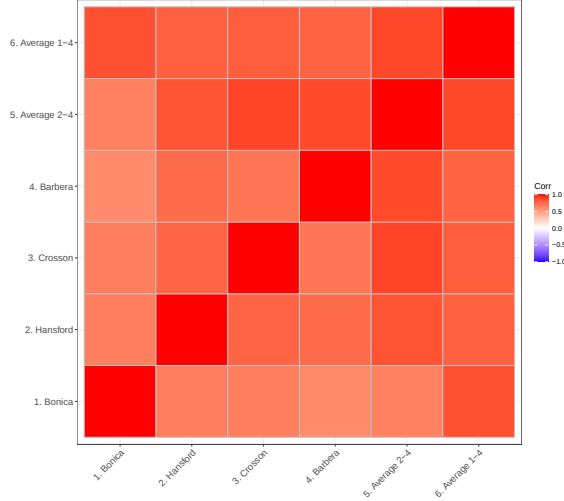


Figure 5: Correlation between ideal points based on existing measures.

Citation	In-Common Orgs	MAD	DIME MAD	Side Agreement	DIME Side Agreement
Crosson et al. (2020)	454.00	0.49	0.70	0.92	0.70
Hansford & Depaoli (2019)	199.00	0.17	0.55	0.90	0.75
Barberá (2013)	200.00	0.81	0.58	0.81	0.74
Average of above	543.00	0.57	0.68	0.83	0.71

Table 5: Cross-validation results, comparing DIME scores to other existing measures.

The scores are correlate well across the measures. Taken together, these results suggest that while there is a strong in-common signal between these measures, they may be capturing different facets of ideology owing to their origins in different interest group behaviors, governed by different strategic considerations in diverse policy domains. Despite the comparability across these scores, we note that using DIME scores produces substantially fewer unscorable organizations than the other three measures or the ensemble measure. As well, only DIME scores include uncertainty estimates that we can use to bootstrap standard errors for our own measure. Finally, DIME scores have been shown to have good validity in a range of settings in a common space, regardless of institutional target (Bonica 2019), though not without criticism (Hill and Huber 2017). For these reasons we continue to prefer DIME scores to seed our measure.

9 Drawbacks and Limitations

While our iterative weighted averaging approach offers key benefits for our application, its simple and nonparametric construction limits it. We discuss a number of these limitations below.

Multidimensionality. Many unsupervised models of ideal point estimation are multidimensional, including Poole and Rosenthal 1985 and Barberá 2015. This multidimensionality is often very useful: identifying the substantive meaning of the second and third dimensions often reveals in-

interesting relationships, and the second dimension is often reported alongside the first. Our method is designed to propagate a set of training labels through a network, and is by design unidimensional. It may be possible to produce multidimensional ideal points with our approach with the following procedure: first, find training labels for a subset of nodes with multiple dimensions. Second, generate ideal points using the first dimension. Third, generate new ideal points using the second dimension, and so on. This approach does not preserve the desirable property that the first and second dimensions are orthogonal, but it may be useful nonetheless.

Time variance. Some unsupervised models like DW-NOMINATE (Poole and Rosenthal 2007) impose structural assumptions in order to estimate smooth, time-varying ideal points, for example by assuming that a legislator’s change in ideal point over time is a random walk with some estimated parameters. Our model, in its simplicity, does not allow for this, instead assuming that organizations’ ideal points are stationary over time. However, it is possible to use our model to estimate time-varying ideal points by estimating it many times using changing subsets of years. For example, we could estimate all organizations’ ideal points using data from 1990-1995, then 1991-1996,...,2016-2021; that 5-year window would be set with the dual constraints of ensuring enough data to estimate ideal points and avoiding overly-long windows in which organizations’ strategic behavior might change. We leave this to future research.

Variance-explained measures. Many unsupervised tools, beginning with principal components analysis, allow for a *variance explained* metric that captures how much of the variation in the adjacency matrix is explained by the first estimated dimension of ideology. The substantive import of this is clear: if the first dimension explains relatively little, it may not be useful for subsequent empirical analysis. Our method does not allow for such analysis, unfortunately. To validate that our measure is useful for empirical analysis we instead rely on face validity and cross-validation.

10 A Data-Generating Process

What data-generating process, and what theory of interest group behavior, might justify our iterative averaging method, and what assumptions underlie that DGP? First we consider the following toy game, then we will explore its assumptions.

10.1 The Game

Consider A organizations, each with a (private) ideal point $p(a)$. When a the Supreme Court grants certiorari to a case, first each organization decides at random whether to sign a brief. Second, Nature selects a Poisson-distributed number of briefs B to sign and places them at random ideal points $p(b)$. Third, each organization that has chosen to sign a brief joins exactly one coalition randomly, where the probability of joining a brief is inversely proportional to the distance between to brief’s ideal point and $|p(a) - p(b)|$. Fourth, after all the coalitions are established, they renegotiate the

briefs' positions such that each brief's ideal point is the unweighted average of the ideal points of the organizations in its coalition.

Under this game, each brief's ideal point is definitionally the average of its signers', and so the cosigning network of a single case may be sufficient to estimate a brief's position. To estimate an organization's position, however, requires many cases.

10.2 The Assumptions

This game makes three key assumptions in its identification, and several others by constrictioin. First, it assumes that firms are not strategic in *whether* they sign a brief. Secondly, it assumes firms are not strategic in *which other firms* they sign with, though they do consider the ideological proximity of the brief. Third, it assumes that in crafting the final brief, each organization has equal ideological weight. Fourth, it assumes that organizations' ideology is constant over time.

Whether to sign. Signing a brief is a costly behavior, both in time and in lawyers' fees. However, organizations benefit from signing briefs in two key ways: first through public policy, and secondly through fundraising and public image (Collins 2018). The empirical evidence for the second motivation is clear, but it is much less clear whether amicus briefs impact judicial outcomes (Spriggs and Wahlbeck 1997; Kearney and Merrill 2000; Box-Steffensmeier, Christenson, and Hitt 2013).

Which firms to sign with. How amicus coalitions form is an opaque and path-dependent process, largely driven by informal networks among interest group leaders and who may have co-signed before. Whereas our model assumes that groups select into coalitions without knowing who else will be part of it, in practice organizations enter coalitions with full knowledge of who is in them and with freedom to remove themselves if they wish.

Equal bargaining power. Next, we assume that each organization in a coalition has equal bargaining power, and therefore a brief's ideal point is the unweighted average of its signers'. It is possible to assume differently: it may be that organizations who have signed more briefs in the past have more expertise and are better at constructing briefs closer to their ideal points. It is likewise possible to perform an iterative *weighted* averaging procedure and estimate ideal points with known weights, but identifying which organizations have more bargaining power *ex ante* is difficult. Measures of network centrality, or the number of signed briefs, or the size of the organization may be good proxies to explore in future work.

Constant ideology. This assumption by construction is weak: there are clear examples of organizations like the National Rifle Association and the AARP shifting ideologically over time. But since these shifts take place over decades rather than months, we believe that it is ignorable for sufficiently few cases.

References

- Barberá, P. 2015. “Birds of the same feather tweet together: Bayesian ideal point estimation using Twitter data.” *Political analysis* 23 (1): 76–91.
- Bonica, A. 2019. “Are donation-based measures of ideology valid predictors of individual-level policy preferences?” *The Journal of Politics* 81 (1): 327–333.
- Box-Steffensmeier, J. M., and D. P. Christenson. 2014. “The Evolution and Formation of Amicus Curiae Networks.” *Social Networks* 36:82–96.
- Box-Steffensmeier, J. M., D. P. Christenson, and M. P. Hitt. 2013. “Quality Over Quantity: Amici Influence and Judicial Decision Making.” *American Political Science Review* 107 (3): 1–15.
- Box-Steffensmeier, J., B. Campbell, D. Christenson, and Z. Navabi. 2018. “Role analysis using the ego-ERGM: A look at environmental interest group coalitions.” *Social Networks* 52:213–227. ISSN: 0378-8733.
- Collins, P. M. 2018. “The Use of Amicus Briefs.” *Annual Review of Law and Science Review* 14:219–237.
- Crosson, J. M., A. C. Furnas, and G. M. Lorenz. 2020. “Polarized Pluralism: Organizational Preferences and Biases in the American Pressure System.” *American Political Science Review* 114 (4): 1117–1137.
- Hansford, T. G., S. Depaoli, and K. S. Canelo. 2016. “Estimating the Ideal Points of Organized Interests in Legal Policy Space.” In *Annual Meeting of the Law and Society Association*, 2–5.
- Hill, S. J., and G. A. Huber. 2017. “Representativeness and motivations of the contemporary donorate: Results from merged survey and administrative records.” *Political Behavior* 39 (1): 3–29.
- Kaufman, A. R., G. King, and M. Komisarchik. 2021. “How to measure legislative district compactness if you only know it when you see it.” *American Journal of Political Science* 65 (3): 533–550.
- Kearney, J. D., and T. W. Merrill. 2000. “The influence of amicus curiae briefs on the Supreme Court.” *University of Pennsylvania Law Review* 148 (3): 743–855.
- Krivitsky, P. N., and M. S. Handcock. 2008. “Fitting position latent cluster models for social networks with latentnet.” *Journal of Statistical Software* 24.
- Li, T., E. Levina, and J. Zhu. 2016. “Network cross-validation by edge sampling.” *arXiv preprint arXiv:1612.04717*.
- Poole, K., and H. Rosenthal. 2007. *Ideology & Congress, 2nd rev. ed.* New Brunswick, NJ: Transaction.
- Poole, K. T., and H. Rosenthal. 1985. “A Spatial Model for Legislative Roll Call Analysis.” *American Journal of Political Science* 29, no. 2 (May): 357–384.
- Spriggs, J. F., and P. J. Wahlbeck. 1997. “Amicus curiae and the role of information at the Supreme Court.” *Political Research Quarterly* 50 (2): 365–386.