

Rational Voter Learning, Issue Alignment, and Polarization*

Martin Vaeth[†]

[Latest version here](#)

October 24, 2025

Abstract

We model electoral competition between two parties when voters can rationally learn about their political positions through flexible information acquisition. Rational voter learning generates polarized and aligned political preferences, even when voters' true positions are unimodally distributed and independent across policy issues. When parties strategically select their positions, voter and party polarization mutually reinforce each other, and both rise as information costs decline. Because we show voters learn exclusively about the axis of party disagreement, party positions respond to only one dimension of aggregate shocks to voter preferences. We then adapt our model to a market setting with horizontally differentiated goods when consumers learn about their product preferences. Lower information costs increase product differentiation and moreover enable firms to charge higher markups, reducing consumer welfare. These results show how lower information costs can reduce welfare in both political and economic contexts.

Keywords: rational inattention, voter ideology, electoral competition, polarization, product differentiation

JEL Classifications: D72, D83, D43, L13

*I am indebted to Roland Bénabou and Alessandro Lizzeri for their guidance throughout the project. For helpful discussions and feedback, I would like to thank Alexander Bloedel, Simon Bornschie, Steven Callander, Mark Dean, Mira Frick, Germán Gieczewski, Francesco Fabbri, Jiadong Gu, Faruk Gul, En Hua Hu, Matias Iaryczower, Ryota Iijima, Gleason Judd, Andreas Kleiner, Botond Kőszegi, Stephan Lauerma, Elliot Lipnowski, John Londregan, Filip Matějka, Nolan McCarty, Dan McGee, Jessica Min, Xiaosheng Mu, Pietro Ortoleva, Jacopo Perego, Wolfgang Pesendorfer, Doron Ravid, Joseph Ruggiero, Fedor Sandomirskiy, Christopher Sims, Guido Tabellini, João Thereze, Leeat Yariv, and seminar participants at Princeton, Bonn, Copenhagen, UCL, Warwick, CERGE-EI, PSE, Berkeley, and the 2025 Conference on Rational Inattention. Financial support was provided by The William S. Dietrich II Economic Theory Center. All errors are my own.

[†]Paris School of Economics. martin.vaeth@psemail.eu

1 Introduction

Voter positions in the United States display two puzzling features. First, one can predict a voter’s position on most policy issues remarkably well by knowing just their location on a one-dimensional left-right axis. As a consequence, positions on different issues are strongly aligned, which is surprising given the wide variety of issues, such as taxation, immigration, and the environment. Second, evidence suggests that voter positions are increasingly polarized, in the sense of being clustered around two poles on the left-right axis. In most settings, we expect the distribution of characteristics to have a unimodal distribution by considerations such as the central limit theorem.¹

This paper provides a joint explanation of issue alignment and polarization based on rational voter learning. Previous research and public discourse attribute issue alignment and polarization to voter biases like confirmation bias, herding due to echo chambers, or partisan news media. By contrast, we show issue alignment and polarized ideology emerge naturally from rational (i.e., no biases) and individual (i.e., no herding) voter learning, driven by voters’ own information choices (i.e., no media effects). Central to the mechanism is that voters learn about their political position through flexible information acquisition in order to decide between two parties. Voter learning can involve understanding the effects of policies—for example, learning about the effects of tariffs to inform their position on trade policy. We assume that such learning is costly, whether in time, effort, or money. This paper shows that the resulting cost-minimization motive structures rational voter learning in a way that generates issue alignment and polarization.

Why study rational voter learning even though voters may be influenced by various biases? An explanation based on rational learning can help disentangle behaviors driven by biases from those arising from rational ideology formation. This distinction is important for assessing the functioning of democratic elections (Achen and Bartels, 2017). If biases dominate, candidates may cater to these biases rather than advancing policies that address societal needs. Conversely, if voter positions reflect rational learning, we may be more optimistic that elections produce policies aligned with voters’ interests. Whether this optimism is warranted is explored in the second part of this paper, which studies the effects of voter learning on party positions.

We analyze rational voter learning in an otherwise standard political-economy model. To study issue alignment, we incorporate a multidimensional policy space, where each dimension represents a different policy issue. Voters face two parties, each adopting a policy platform in this space. These policy platforms are exogenous in the first part of our paper. Following the literature, a voter’s utility decreases quadratically in the distance between a policy and her *ideal point*, which reflects her political position. In our model, voters initially lack knowledge of their ideal point but

¹Section 3.4 examines the evidence for issue alignment and ideological polarization in detail. While the evidence for issue alignment is strong, the evidence on bimodality is more speculative; there is ongoing debate about this issue (Fowler, Hill, Lewis, Tausanovitch, Vavreck, and Warshaw, 2022). Understanding the sources of potential polarization is important because of evidence of increasing voter polarization (Pew Research Center, 2014; Martin and Yurukoglu, 2017; Eguia and Hu, 2022) and because of the relevance of polarization for political tensions, as emphasized by the literature on polarization (Esteban and Ray, 1994): a polarized distribution can lead to the formation of two homogeneous groups with few individuals bridging the divide, increasing the risk of conflict.

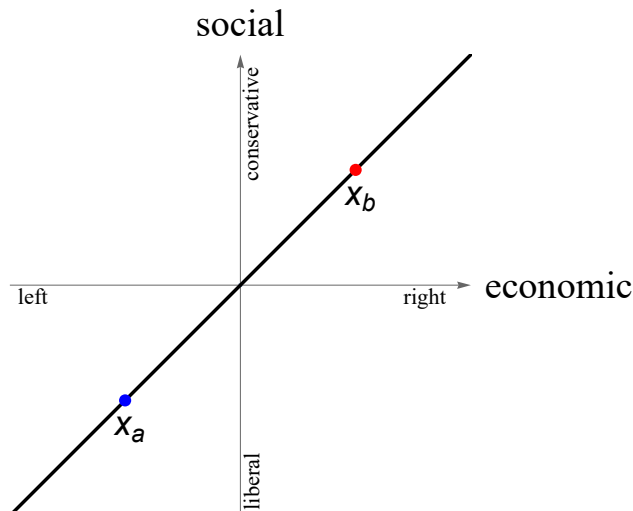


Figure 1: Exemplary policy platforms x_a and x_b of parties a and b , respectively, under two policy issues. Voter learning induces revealed ideal points on the diagonal line.

can form political opinions through learning. We assume that voters start with a homogeneous prior that conforms to the true distribution of ideal points, which in the simplest case is normal and independent across issues.² Voters then update their beliefs by acquiring information flexibly, following the rational-inattention framework (Sims, 2003). This framework allows voters to choose both how much and what type of information to acquire, subject to a cost increasing in information. After acquiring information, and given quadratic utility, a voter selects the party whose policy bundle is closer to her expected ideal point. We refer to this expected ideal point as her *revealed ideal point*, distinguishing revealed ideology—the distribution of revealed ideal points induced by learning—from the distribution of true ideal points.

Our first result is that rational voter learning generates issue alignment and polarization: revealed ideal points align along a left-right axis and cluster around two poles, even if true ideal points are independent across issues and centrally clustered. At the core of this mechanism is that voters trade off making an informed choice and minimizing the cost of information. By acquiring only the information necessary to determine which party position is closer, voters align their preferences along a single axis and polarize—a process we explain in more detail next.

To understand the high-level intuition for why cost-effective learning generates issue alignment, consider a two-dimensional policy space with an economic and a social issue. Suppose the two parties propose policy platforms, x_a and x_b , as shown in Figure 1: one platform is more economically left and socially liberal than the other. A unique line passes through these platforms, the direction of which we refer to as the *axis of disagreement*. This axis of disagreement is the direction along which parties differentiate, combining the two issues in proportion to the extent of disagreement on each. Crucially, voters learn only where their ideal point lies along this axis, as this determines which platform is closer. Their position along directions orthogonal to the axis is irrelevant, as

²More generally, the paper allows for any elliptical distribution with arbitrary correlation across issues. Section 3.1 discusses the case of heterogeneous priors.

it does not affect the *relative* proximity of the platforms. Consequently, voters align preferences across issues, learning only whether they lean economically left and socially liberal or economically right and socially conservative. Revealed ideal points lie on a line, which can be interpreted as an endogenous left-right axis. This left-right axis reflects the direction of party disagreement and passes through voters' prior expectation of their ideal points. This result holds for any number of policy issues, assuming a reflection-invariant information cost and an elliptical prior.

The mechanism by which cost-effective learning generates ideological polarization is as follows. Voters only need to determine which party they are closer to, not by how much. To do so, they optimally acquire a *binary* signal about their ideal points. This signal sorts voters into two groups: those who believe they are closer to one party and those who believe they are closer to the other. Even centrist voters categorize themselves based on this binary information, causing them to perceive their preferences as more extreme than they truly are. We also examine how robust this mechanism is to shocks about the desirability of the two parties that occur after learning. For example, voters may receive new information about candidates' competence closer to the election. These "valence shocks" make it valuable for voters to learn *how much* they prefer one party over the other, which previously held no value. We show that the distribution of revealed ideal points remains confined to the endogenous left-right axis. For small valence shocks, this distribution is bimodal.

The second part of this paper examines whether elections produce policies that serve voters' interests when voter positions result from rational learning. To address this, we endogenize party platforms and explore their interaction with voter learning under two alternative timings. In both timings, parties strategically position themselves in the policy space, balancing electoral prospects against their own policy preferences. In the first timing, voters learn before parties choose their platforms; in the second, parties move first and voter learning responds to party positioning. Considering both timings allows us to compare how the order of information acquisition and platform choice shapes equilibrium outcomes.

In the first timing, voters learn optimally in anticipation of party platforms, and parties choose their platforms based on the revealed voter ideology resulting by this learning. We show two results about party platforms: more voter learning can hurt voters by increasing party polarization, and parties respond to only one dimension of aggregate shocks to voter preferences.

First, ideological polarization of voters and platform polarization are mutually reinforcing, and perhaps counterintuitively, less costly information increases both. When party platforms are more polarized—that is, farther apart—voters face higher stakes in the election and are motivated to learn more about their ideal points. More information leads to a more polarized distribution of voter ideal points, creating more extreme voters who are less responsive to party platforms. As a result, parties can move their platforms closer to their ideal policies without losing as many votes, reinforcing platform polarization. Cheaper information amplifies this cycle: it encourages voters to learn more about their preferences, leading to more voter and platform polarization. Paradoxically, better access to information about political preferences may harm voters, as the

equilibrium platforms end up farther from the welfare-maximizing policy. This mechanism may help explain the increasing platform polarization in the US over recent decades (McCarty, Poole, and Rosenthal, 2016), during advances in information technology, such as the internet, which has made information more accessible.

Second, due to rational voter learning, policy platforms respond solely to a single dimension of aggregate shocks to voter preferences. Aggregate shocks to voter preferences affect what the optimal policy should be across multiple dimensions; ideally, we would want voters to learn about these shocks so that parties can adjust their platforms accordingly. However, voters' optimal learning reduces politics to a single dimension—the axis of party disagreement—and neglects all other dimensions. As a result, party platforms respond solely to one dimension of aggregate shocks. This inefficiency is particularly problematic because it is not apparent from revealed voter ideology: aggregate shocks manifest as one-dimensional changes along the axis of disagreement, and party platforms adjust in response to these observed changes. Meanwhile, the other dimensions of aggregate shocks remain latent—they do not influence revealed voter ideology and thus go unnoticed in empirical data.

In the second timing, parties choose their platforms before voters learn. This gives parties an agenda-setting role: through their policies, parties influence which issues voters pay attention to. For example, if a party polarizes on a policy issue, voters will pay more attention to this issue as it becomes more relevant to the electoral choice. This timing introduces two novel strategic forces: a moderation force and a differentiation force. First, parties may moderate their policy platforms to skew voter learning in their favor. Second, the more extreme party may find it optimal to differentiate from the moderate party to trigger more voter learning, thereby reducing the skew toward their opponent. As information costs decline, the moderation force weakens, leading to increased platform polarization, as in the first timing. The differentiation force implies that parties may adopt positions more extreme than their own ideal policies to strategically affect voter learning. In contrast, in models with exogenous voter positions (e.g., Roemer, 1997), parties never adopt positions more extreme than their own ideal policies.

To illustrate the broader applicability of our results, we adapt our model to a market setting where consumers learn about their preferences for horizontally differentiated products. In this context, firms not only choose product attributes—similar to how parties select policy platforms—but also set prices to maximize profits. We show that cheaper information leads consumers to become better informed, prompting firms to increase product differentiation. While this could benefit consumers by better matching them to products, product differentiation harms consumers because firms exploit it to raise prices. As a result, despite the reduction in information costs, consumer welfare decreases overall.

Our results call for caution when interpreting empirical findings about voter ideology being influenced by party elites as a sign of voter irrationality. Political scientists have long argued that political elites have a large influence on the ideology of voters (Campbell, Converse, Miller, and Stokes, 1960; Zaller, 1992; Lenz, 2012). Such findings have traditionally been interpreted as

evidence of voter irrationality. This interpretation is shared by Achen and Bartels (2017), who argue such voter behavior presents a serious threat to democracy. If parties can shape the ideology of voters instead of merely responding to it, it is unclear whether elections produce governments responsive to the preferences of voters. We hope to contribute to this debate by showing that some forms of party influence on voter ideology are consistent with voter rationality and do not preclude that policy is responsive to voters’ true preferences. In our model, both issue alignment and polarization of voter ideal points depend on party platforms. First, the alignment of voter ideal points across policy issues is determined by the relative positions of parties. As illustrated in Figure 1, because party a is more left and liberal than party b , in the resulting voter ideology a left economic position aligns with a liberal social position. Second, more polarized party platforms result in more polarized voters (Proposition 1). Although in both cases voters seem to simply follow party positions, our model shows such effects result from rational voter learning. Moreover, in our equilibrium, party platforms do respond to voters’ true ideal points, namely to the center of their distribution (Theorem 3). However, parties do not respond to aggregate shocks to voter preferences in more than one dimension (Theorem 4). On a higher level, our model illustrates how revealed preferences may differ systematically from true preferences and how they may do so in a context-dependent way.

Our model has policy implications for addressing issue alignment and polarization. Traditional approaches—such as improving political knowledge (Carpini and Keeter, 1996) or breaking up echo chambers (Sunstein, 2018)—may be ineffective if issue alignment and polarization stem from rational voter learning. Instead, more fundamental changes to voters’ choice sets would be necessary to incentivize multidimensional and non-binary learning. Multidimensional learning could be incentivized through electoral reforms that increase the number of parties, such as transitioning from plurality elections to proportional representation (as suggested by Corollary 1).³ Non-binary learning could be incentivized through voting or participation mechanisms that elicit the intensity of voter preferences (e.g., Casella, 2005).

The paper is organized as follows. The rest of this section discusses related literature. Section 2 introduces the model. Section 3 analyzes voter learning and discusses the related evidence. Section 4 studies electoral competition when voter ideology results from optimal learning. Section 5 studies the model under an alternative timing, in which parties move before voters learn. Section 6 adapts our model to an industrial organization setting with horizontally differentiated products. Section 7 concludes.

Related Literature This paper contributes to the growing literature on rationally inattentive voters. We show rational inattention explains properties of voter ideology by studying *flexible* learning about *ideal points*. Matějka and Tabellini (2021) study electoral competition where voters are inattentive to *party platforms*. They show more attentive voters influence platforms more strongly, as they respond more to them. By contrast, we show attention to *ideal points* reduces respon-

³Under $k > 2$ parties, a weaker form of issue alignment is predicted by our model: revealed voter ideology is at most $(k - 1)$ -dimensional.

siveness to platforms, which increases polarization when information becomes cheaper. Matějka and Tabellini (2021) find that multidimensional policies are inefficient because voters pay excessive attention to divisive issues. In our model, inefficiency arises because voters focus on a *single* dimension, the axis of party disagreement. Yuksel (2022) analyzes voters learning under partitional signals and finds that more specialized learning increases polarization. We allow for flexible learning and show that cheaper information increases polarization through a different mechanism. Li and Hu (2023) study attention to implemented policies in an electoral-accountability setting. They show that the welfare effects of increased attention and mass polarization are ambiguous. Hu, Li, and Segal (2023) analyze learning about valence through an attention-maximizing intermediary. They show it generates policy divergence even with office-motivated candidates. By contrast, we study divergence as arising from ideologically motivated candidates and study how it interacts with voter learning.

The literature has proposed several explanations for issue alignment and polarized ideology. Converse (1964) introduced the concept of issue alignment, which he termed ideological constraint, attributing it to logical, psychological, and social sources. Enke, Rodríguez-Padilla, and Zimmermann (2023) propose that moral universalism, the degree of altruism toward strangers versus in-group members, explains correlations in policy views across domains. Spector (2000) derives one-dimensional ideology from cheap talk between two groups with differing priors, while DeMarzo, Vayanos, and Zwiebel (2003) attribute it to networks and a persuasion bias. Bayesian persuasion can also generate low-dimensional types (Rayo and Segal, 2010; Malamud and Schrimpf, 2022).⁴ Ideological polarization has been linked to cognitive limitations, such as correlation neglect (Levy and Razin, 2015; Ortoleva and Snowberg, 2015) and bounded rationality (Eguia and Hu, 2022). Other papers have also proposed rational inattention as a mechanism for belief polarization, albeit through ex-ante heterogeneity in preferences (Novák, Matveenko, and Ravaioli, 2024) or path-dependent sequential information acquisition (Nimark and Sundaresan, 2019). Callander and Carbajal (2022) explain dynamic polarization through voters adjusting their ideal points toward the party they voted for to rationalize their choice. Our model complements these works by providing a unified explanation for both issue alignment and polarization.

The literature has provided many mechanisms for platform divergence, breaking the median voter result by Downs (1957). We do not propose a new mechanism but show how the new ingredient of our model—endogenous ideology formation through voter learning—interacts with perhaps the most prominent mechanism for platform divergence: ideologically motivated parties and probabilistic voting. Ideological parties still converge to the median voter unless the electoral outcome is uncertain (Wittman, 1983; Hansson and Stuart, 1984; Calvert, 1985). The two common ways to introduce electoral uncertainty are through uncertainty about the ideological position of the median voter (Roemer, 1994) and through valence shocks (Hinich, 1977; Lindbeck and Weibull, 1987). Our model falls into the latter category, which, according to Duggan (2017), has seen little,

⁴More broadly related is McMurray (2023), who shows that pivotality considerations in multidimensional common-value elections drive party platforms to bundle logically related issues, reducing platforms to a persistent one-dimensional axis.

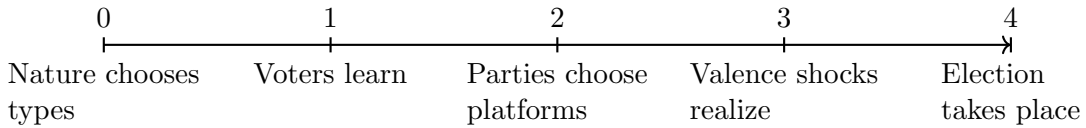
if any, formal analysis under ideologically motivated parties. We introduce a version of this model that is tractable even in a multidimensional policy space. The model generalizes the mean-voter theorem by Hinich (1977) and allows comparative statics with regards to platform polarization.

Finally, we contribute to the burgeoning theoretical literature on rationally inattentive consumers in industrial organization. As we illustrate in section 6, our model can be adapted to a setting of consumer learning in the face of horizontally differentiated products. Of particular relevance are Albrecht and Whitmeyer (2023) and Biglaiser, Gu, and Li (2024), who study a duopoly with consumers learning about their preferences. Albrecht and Whitmeyer (2023) show consumers learn only about the relative value of products and that, in contrast to Ravid, Roesler, and Szentes (2022), an ex-post efficient equilibrium exists as the information cost converges to zero. Biglaiser, Gu, and Li (2024) study comparative statics of the unique symmetric equilibrium and show an application to platform design. In both papers, as is standard, consumers learn directly about their valuations of products, whereas we assume they learn about their ideal products in an attribute space. The additional structure on preferences facilitates an analysis of endogenous product *attributes*, whereas the literature typically focuses on the effect of attention on firm’s *pricing* decisions. An exception to this is Cunha, Osório, and Ribeiro (2022), who study a spatial setting where consumers pay attention to product attributes and prices, whereas in our model consumers learn about their preferences. As a result, our model makes predictions on the distribution of consumer preferences, which are endogenously one-dimensional allowing us to tractably study a multidimensional attribute space.

2 Model

We employ a standard probabilistic voting model with valence shocks (Hinich, 1977; Lindbeck and Weibull, 1987) and add to it an earlier stage in which voters learn about their ideal points, anticipating the election. We discuss our assumptions at the end of this section.

Game The policy space is $\mathbb{R}^n$ with $n \in \mathbb{N}$, where each dimension corresponds to a policy issue. Voter ideal points as well as platforms live in this space. There is a continuum of voters $i \in [0, 1]$ and two parties, a and b . The timing is as follows.



- (0) Nature draws voters’ ideal points $\theta_i \in \mathbb{R}^n$ independently from an elliptical distribution $\mu \in \Delta(\mathbb{R}^n)$, with mean normalized to 0 and an arbitrary positive definite covariance matrix $\Sigma \in \mathbb{R}^{n \times n}$.⁵

⁵For any topological space X , we denote by $\Delta(X)$ the set of Borel probability measures on X . A measure $\mu \in \Delta(\mathbb{R}^n)$ with mean 0 and covariance matrix Σ is *elliptical* if its characteristic function Φ takes the form $\Phi(\theta) = \psi(\theta^\top \Sigma \theta)$ with $\psi: \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$. If μ admits a density f , it must be of the form $f(\theta) = g(\theta^\top \Sigma^{-1} \theta)$ with $g: \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$. This

An example is the normal distribution $\mathcal{N}(0, \Sigma)$, but elliptical distributions can also have bounded support.

- (1) Each voter $i \in [0, 1]$ chooses how to learn about their ideal point $\theta_i \in \mathbb{R}^n$.
- (2) Parties a and b observe the realized distribution of voter preferences and commit to policy platforms $x_a \in \mathbb{R}^n$ and $x_b \in \mathbb{R}^n$, respectively.
- (3) Each voter observes party platforms (x_a, x_b) and the realization of the valence shocks $\nu_i \in \mathbb{R}$, which can be interpreted as a signal about the competence differential of party candidates.
- (4) Voters elect their preferred parties. The party with the majority wins, and each party wins with probability one half if a tie occurs.

Because each voter is infinitesimal and there is no aggregate uncertainty about voter ideal points, this timing is equivalent to one where stages one and two occur simultaneously. If, alternatively, parties moved before voters learned, the results on voter learning (Theorem 1 and 2) would remain unchanged. The qualitative predictions on voter ideology do not depend on whether voters learn given observed or given anticipated platforms. Our timing rules out that parties choose platforms in order to affect voter learning. In section 5, we show the comparative statics on polarization (Theorem 3) extend to the reversed timing.

Voters The utility U_i of voter $i \in [0, 1]$ has three components, which we expand on below. Utility depends on the implemented policy x and the voter’s ideal point θ_i via the policy utility $u(x, \theta_i)$, on the net valence shock ν_i for candidate b , and on the information cost $c(\tau_i)$ of signal structure τ_i scaled by κ ,

$$U_i(x, \tau_i) = u(x, \theta_i) + \nu_i \cdot \mathbb{1}_{\{\text{vote } b\}} - \kappa \cdot c(\tau_i). \quad (1)$$

Following the workhorse model of spatial voting, we assume voters’ policy utility is quadratic in the difference between their ideal point and the policy; that is,

$$u(x, \theta) = -(x - \theta)^\top A(x - \theta), \quad (2)$$

where $A \in \mathbb{R}^{n \times n}$ is an arbitrary symmetric, positive definite matrix. Although this assumption is restrictive, most of the evidence on voter ideal points assumes quadratic utility with a homogeneous A and shows it can explain voters’ survey responses well (see section 3.4).⁶ Further, quadratic utility allows us to speak about “revealed ideal points” of voters who have a non-degenerate belief π about their true ideal point θ , as the following remark shows.

Remark 1. *A voter with belief π over her ideal point θ votes as if she had the known ideal point $\mathbb{E}_\pi[\theta]$.*

symmetry assumption states that the isodensity curves are ellipses. Independence of ideal points is not necessary for our results on voter ideology in section 3.

⁶Moreover, in recent work, Bachmann, Sarasua, and Bernstein (2024) find that out of a range of commonly used algorithms, the one based on quadratic utility performs best at predicting survey responses.

By a bias-variance decomposition of the expected utility from policy x under belief π over θ ,

$$\mathbb{E}_\pi[u(x, \theta)] = u(x, \mathbb{E}_\pi[\theta]) - \mathbb{E}_\pi[(\theta - \mathbb{E}_\pi[\theta])^\top A(\theta - \mathbb{E}_\pi[\theta])]. \quad (3)$$

The latter term does not depend on x , so when a voter compares two platforms (or survey response items), they choose the one that is closer to their posterior mean. Hence, a voter with belief π acts like a voter with a known ideal point of $\mathbb{E}_\pi[\theta]$. Accordingly, we call the posterior mean of a voter's belief π over θ her *revealed ideal point*. This ideal point is the one that is estimated from survey responses, which is important when we interpret empirical findings about voter ideology. We refer to the distribution $\rho \in \Delta(\mathbb{R}^n)$ over posterior means induced by learning as the *revealed ideology* in the population, to distinguish it from the true distribution over ideal points.

The valence shock ν_i is to be interpreted as the valence difference between parties b and a . It has the same distribution for all voters $i \in [0, 1]$. The distribution of ν has a finite first absolute moment and admits a continuous density f_ν that is symmetric around 0, and strictly quasiconcave. The symmetry of the valence shocks means that no party has a valence advantage, which simplifies our analysis of electoral competition. Strict quasiconcavity together with symmetry implies the density of the valence shock is maximal at 0. We show later that this assumption implies more extreme voters are less sensitive to party platforms. Because we assume parties care only about their expected vote shares, we do not need to specify the joint distribution of valence shocks.

Learning Voters share a homogeneous prior μ , conforming to the true distribution, before learning.⁷ Each voter can acquire *any* signal structure (Blackwell experiment) about her ideal point at a cost proportional to mutual information, as in the rational-inattention literature (Sims, 2003; see also the survey Maćkowiak, Matějka, and Wiederholt, 2023). The information cost captures that learning takes time and effort. To define mutual information, recall that a signal structure specifies a conditional distribution over signal realizations given any ideal point. Upon a signal realization, the agent forms a posterior via Bayesian updating. Thus, a signal structure induces a distribution over posteriors. Bayesian updating implies this distribution averages to the prior, also called Bayes consistency. In fact, following the posterior approach (Kamenica and Gentzkow, 2011; Caplin and Dean, 2013), we can represent signal structures as Bayes-consistent distributions $\tau \in \Delta(\Delta(\mathbb{R}^n))$ over posteriors $\pi \in \Delta(\mathbb{R}^n)$.⁸ The mutual-information cost can then be defined as the expected Kullback-Leibler divergence⁹ of posterior π from prior μ ,

$$c(\tau) = \mathbb{E}_\tau[D_{\text{KL}}(\pi||\mu)]. \quad (4)$$

⁷In section 3.1, we show Theorem 1 can be extended to allow for ex-ante heterogeneity in beliefs.

⁸That Bayesian updating imposes only Bayes-consistency of τ holds for general Polish state spaces, which includes $\mathbb{R}^n$, as a consequence of the disintegration theorem, as shown by Lipnowski and Ravid (2023), Appendix C.2.

⁹The Kullback-Leibler divergence of π from μ is defined as

$$D_{\text{KL}}(\pi||\mu) = \begin{cases} \int_{\mathbb{R}^n} \log\left(\frac{d\pi}{d\mu}\right) d\pi & \text{if } \pi \ll \mu \\ \infty & \text{else,} \end{cases}$$

where $\frac{d\pi}{d\mu}$ is the Radon-Nikodym derivative and $\pi \ll \mu$ means π is absolutely continuous with respect to μ .

Intuitively, the Kullback-Leibler divergence defines a “distance” on beliefs, and mutual information measures how much the acquired information moves the voter’s belief away, on average, from her prior according to this “distance.” We assume different voters’ signal realizations are independent. The cost parameter κ in (1), which we vary for comparative statics, translates mutual information into utils.

For voters to acquire costly information despite never being pivotal among the continuum of voters, we assume voters engage in *expressive voting*, as is standard in the literature on rationally inattentive voters (Matějka and Tabellini, 2021; Hu, Li, and Segal, 2023; Li and Hu, 2023).¹⁰ That is, voters genuinely care about voting for the correct candidate given their true preferences, for which they are willing to incur an information cost. The reason may be that voters derive a psychological benefit from doing so or they consider it their civic duty (see also Feddersen and Sandroni, 2006).

Formally, voter i first chooses information τ_i and, after the observation of valence ν_i and platforms (x_a, x_b) , votes for $x \in \{x_a, x_b\}$ to maximize $U_i(x, \tau_i)$. Dropping indices, the voter’s choice of information, that is, distribution $\tau \in \Delta(\Delta(\mathbb{R}^n))$ over posteriors $\pi \in \Delta(\mathbb{R}^n)$ that is Bayes-consistent (BC), must solve the following problem:

$$\sup_{\tau \in \Delta(\Delta(\mathbb{R}^n))} \int \left(\mathbb{E}_\nu \left[\max \left\{ \mathbb{E}_\pi[u(x_a, \theta)], \mathbb{E}_\pi[u(x_b, \theta)] + \nu \right\} \right] - \kappa D_{\text{KL}}(\pi || \mu) \right) d\tau \quad (\text{P})$$

$$\text{s.t. } \int \pi d\tau = \mu. \quad (\text{BC})$$

The integrand of (P), which we call the *value function*, has the following interpretation. Given a posterior π , the voter anticipates that for each realization of the valence shock ν , they will choose the maximum out of the expected policy utility of party a , $\mathbb{E}_\pi[u(x_a, \theta)]$, and the expected policy and valence utility of party b , $\mathbb{E}_\pi[u(x_b, \theta)] + \nu$. Further, they incur a cost proportional to the Kullback-Leibler divergence $D_{\text{KL}}(\pi || \mu)$. In Appendix D.2, we establish that an optimal distribution τ over posteriors exists, despite the infinite and non-compact state space.

The distribution τ over posteriors induces a distribution ρ over posterior means, which are well-defined by existence of the prior mean. Because voters are ex-ante homogeneous, we assume all voters acquire the same information τ .¹¹ Then, given voters’ ideal points and signal realizations are uncorrelated, the population distribution of revealed ideal points equals ρ (Uhlig, 1996).¹²

¹⁰Martinelli (2006) studies information acquisition in large elections assuming the pivotal-voter model. In large electorates, all voters are nearly uninformed.

¹¹We expect this assumption to be without loss. Even if multiple optimal τ ’s existed and different voters acquired different ones, the resulting population distribution of revealed ideal points should be equivalent to one where each voter chooses the population-mixture of τ ’s. Such a choice of information τ is also optimal because the posterior-separable information cost implies indifference to mixing between optima.

¹²The distribution ρ is necessarily a mean-preserving contraction of the prior, which has finite second moments, so ρ has finite second moments. Thus, the law of large numbers by Uhlig (1996) applies if we interpret the population distribution ρ as a Pettis integral.

Parties Two parties, a and b , choose platforms, x_a and x_b , respectively, to maximize a weighted sum of their expected vote share and their ideological utility.¹³ Their utilities, U_a and U_b , as a function of platforms, x_a and x_b , and the population distribution of revealed ideal points $\rho \in \mathbb{R}^n$, are

$$U_a(x_a, x_b, \rho) = m \int_{\mathbb{R}^n} F_\nu(u(x_a, \theta) - u(x_b, \theta)) d\rho(\theta) + u(x_a, x_a^*) \quad (5)$$

$$U_b(x_a, x_b, \rho) = m \left(1 - \int_{\mathbb{R}^n} F_\nu(u(x_a, \theta) - u(x_b, \theta)) d\rho(\theta) \right) + u(x_b, x_b^*) \quad (6)$$

where $m > 0$ is the weight on vote share and x_j^* is the known ideal point of party $j \in \{a, b\}$. The probability of voting for a given revealed ideal point θ is the probability that the valence shock ν does not exceed $u(x_a, \theta) - u(x_b, \theta)$, that is, $F_\nu(u(x_a, \theta) - u(x_b, \theta))$. The expected vote share is simply this vote probability integrated over all voters. We assume the parties have different ideal points, $x_a^* \neq x_b^*$, which guarantees platform divergence in equilibrium. Otherwise, voters would have no incentive to learn, resulting in a trivial equilibrium.

Equilibrium We study pure-strategy perfect Bayesian equilibria. In the last period, voters vote for their preferred platform given their revealed ideal point and the realized valence shock. Before that, parties simultaneously choose platforms x_a and x_b given the distribution of revealed ideology ρ induced by voter learning. Voters, in turn, learn optimally anticipating platforms x_a and x_b . Equilibria can be characterized by a triple (ρ, x_a, x_b) , where ρ is the distribution over posterior means induced by a solution τ to (P) given (x_a, x_b) , x_a maximizes (5) given (ρ, x_b) , and x_b maximizes (6) given (ρ, x_a) . Intuitively, revealed voter ideology ρ results from optimal voter learning given the anticipated platforms (x_a, x_b) , which optimally respond to each other and to revealed voter ideology.

2.1 Discussion

Learning about Ideal Points Voter learning about ideal points can be interpreted as (i) introspecting on how to *value* the consequences of policies, (ii) learning about the private *consequences* of policies (recall ideal points are private), or (iii) a combination of both. As an example, consider a voter's position on income taxation. To determine her optimal tax policy, the voter may want to introspect on her values for equity versus efficiency, as well as learn about what tax bracket she is in and what other economic consequences the policy has. We remain agnostic as to which interpretation should be adopted.

We assume voters can acquire costly information about their ideal points, but they observe party platforms and a valence signal for free. This approach allows us to make clear what mechanisms result from endogenous voter learning about ideal points, as opposed to learning about platforms or valence (for the latter, see Matějka and Tabellini, 2021; Hu, Li, and Segal, 2023). However, as

¹³That only two relevant parties exist is typically understood as a consequence of plurality voting systems (Duverger, 1954). For multiple parties, see also Corollary 1 and the discussion after Theorem 2.

we briefly illustrate in Appendix D.1, our results on voter ideology also hold when voters do not know platforms but observe a common signal about platforms, based on which they choose how to learn about ideal points. What is more importantly ruled out by our assumption is that voters learn *jointly* about platforms and ideal points, which is an interesting avenue for future research.

Flexible Information Acquisition The rational-inattention approach preserves tractability while allowing complete flexibility in what kind of information voters can acquire. The flexibility assumption ensures the optimal signal structure is determined endogenously and not through exogenous restrictions. In particular, we are not imposing that signals about different policy issues need to be independent. That is, voters can, for example, acquire a signal that informs them about whether they are left or right, when aggregating their positions on multiple policy issues. We show such signals are, in fact, optimal.

The substantive meaning of this assumption depends on which of the above-mentioned interpretations of voter learning we adopt. When we interpret voter learning as learning about values, one can think of the voter imagining two policies that differ on multiple issues and introspecting on their relative desirability, similar to drift-diffusion models, widely used in psychology and neuroscience. By not imposing any restrictions on information, our approach stays true to the original motivation of rational inattention as modeling the brain as an efficient information processor subject to only information capacity constraints (Sims, 2003). On the other hand, when we interpret voter learning as learning about private consequences of policies, another way for such learning to be aggregated across dimensions is through information intermediaries, as in Hu, Li, and Segal (2023). Voters may learn about private policy consequences from sufficiently personalized media outlets, such as news feeds or newspapers catering to specific demographics. Such media outlets may aggregate information about different policy issues to a one-dimensional signal, as other models of media assume (Duggan and Martinelli, 2011; Yuksel, 2022; Perego and Yuksel, 2022).

Mutual Information Cost Although we assume the standard mutual-information cost, our results hold more generally. We use only posterior separability, Blackwell monotonicity, (reflection)-invariance, and continuity properties of the information cost. A posterior separable cost (Caplin, Dean, and Leahy, 2022) is linear under mixing between distributions over posteriors, which we use in our proof of Theorem 1. Posterior separability has foundations from information theory (Sims, 2003), sequential sampling (Morris and Strack, 2019; Bloedel and Zhong, 2020; Hébert and Woodford, 2023), and constant marginal cost of experimentation (Pomatto, Strack, and Tamuz, 2023). Blackwell monotonicity means less information, in the sense of a garbling, is less costly. This property implies agents will not acquire information that does not affect their behavior, because ignoring such information would be cheaper. Reflection-invariance of the Kullback-Leibler divergence is used for Theorem 1. In Appendix D.3, we discuss this property further and show it is satisfied by certain versions of distance-based information costs that the literature has recently proposed. Our results assuming a normal distribution (Proposition 1 and Theorem 3), use the stronger property of invariance stemming from information geometry (Amari, 2016; Caplin, Dean, and Leahy, 2022). Finally, we use lower semicontinuity of the Kullback-Leibler divergence to

establish existence and continuity results.

Party Objective Our party objective makes two notable assumptions. First, as is common to the probabilistic voting literature, we assume parties care about their expected vote share instead of the probability of winning (for examples, see the references in Duggan, 2017). The expected vote share is in general a less complex object and, under some conditions, equivalent to the probability of winning.¹⁴ Second, we model ideological motivation through an additively separable objective. The more common approach, following Wittman (1973), assumes parties care about the implemented policy. Our objective captures in a simpler way a party (or party candidate) that cares both about votes and about not deviating too far from the parties’ ideology. The advantage of our party objective is that it provides greater tractability—see our discussion under related literature—while capturing the main trade-off between vote share and ideology.

3 Voter Learning

We characterize optimal voter learning given equilibrium party platforms, x_a and x_b , assuming $x_a \neq x_b$ (otherwise, voters will learn nothing). In section 4 on electoral competition, we show parties indeed choose distinct platforms in equilibrium if the party ideal points are distinct.

This section can be seen independent of the political-economy application, and results apply analogously to an industrial organization setting with horizontally differentiated goods in a product attribute space. In that setting, valence shocks can be seen as uncertainty about prices, for example.

3.1 Issue Alignment

Our first result shows the revealed ideal points of voters (their posterior means) are on a line. Thus, even though the true distribution of ideal points is multidimensional, the revealed ideology in the population is one-dimensional. By implication, the revealed ideal points are perfectly aligned across policy issues, which holds even if the true ideal points are independent across dimensions. We show in section 3.4 the data on voter ideal points indicates that ideal points are on a line (Proposition 2). All proofs are relegated to the Appendix.

Theorem 1 (Issue Alignment). *The distribution of revealed ideal points ρ has support inside the line through the prior mean with direction $\Sigma A(x_b - x_a)$.*

On a high level, the intuition of this result is that only one dimension of the ideal point is relevant for voting. We outline the logic of the proof of Theorem 1 more carefully for the special case $A = \Sigma = I_n$, namely, that the matrix A associated with the policy utility and the prior covariance matrix Σ are equal to the identity matrix. In this case, the line of voter ideal points is parallel to the platform difference $x_b - x_a$, as in Figure 1. The first part of the proof shows that under quadratic

¹⁴Patty (2002) and Patty (2005) provide conditions for equivalence between maximizing probability of winning and expected vote share under office-motivated candidates. Yuksel (2022) gives a condition under which probability of winning equals the expected vote share under ideologically motivated candidates. More generally, one could assume parties care non-linearly about their expected vote share. We expect Theorem 3 to be robust to this extension.

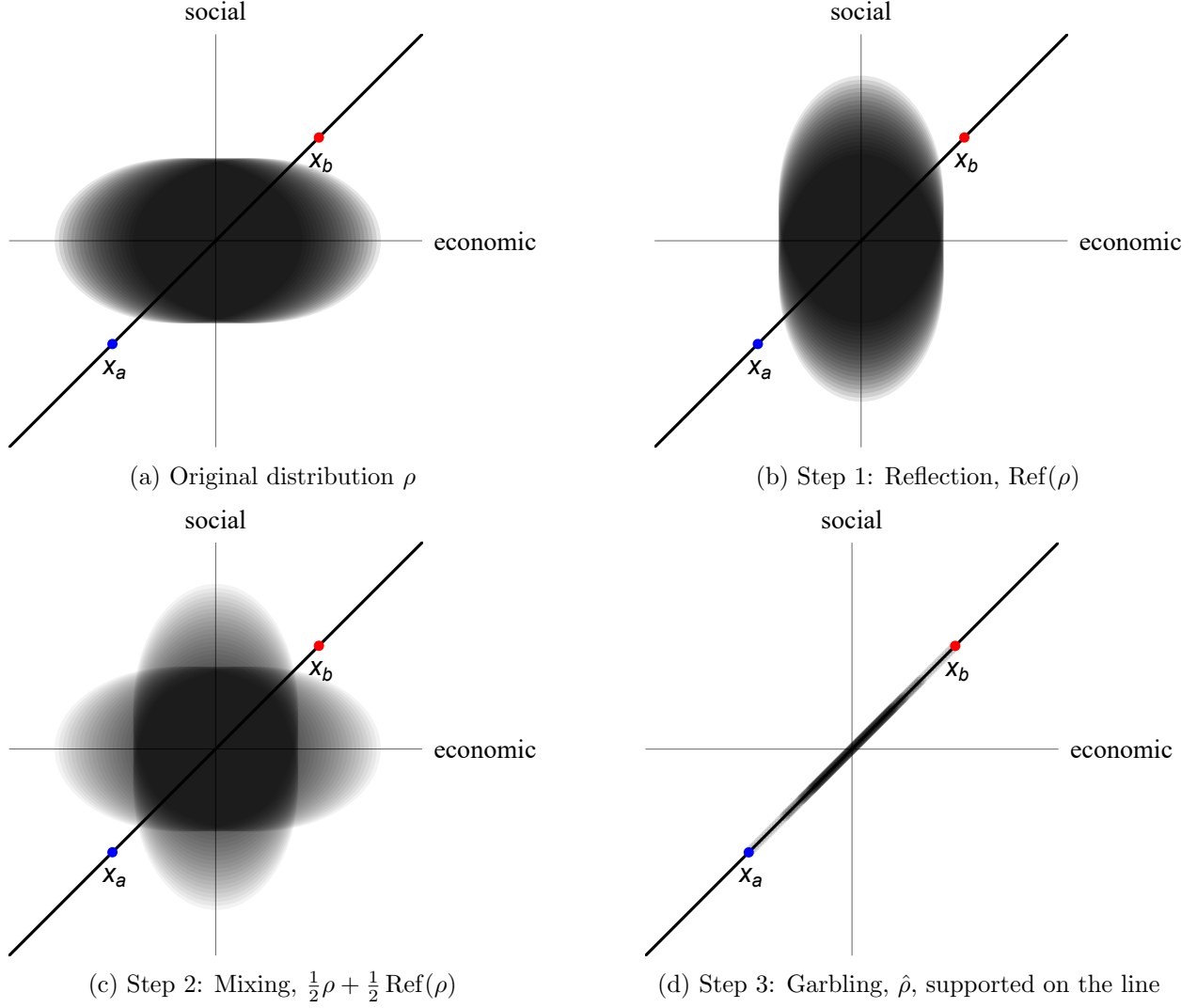


Figure 2: Reflection argument underlying the proof of Theorem 1. The dark clouds visualize the distribution over posterior means.

utility, the instrumental value of information depends only on the projection of the posterior mean on the platform difference $x_b - x_a$.¹⁵ For the second part, suppose by way of contradiction that voters acquired some information such that the induced distribution ρ over posterior means was not supported on the diagonal line in Figure 2, (a). The proof constructs through a reflection argument in three steps a distribution over posteriors that has the same instrumental value but that is cheaper. For the first step, the voter is indifferent between the original information and acquiring the “reflected” distribution over posterior means in Figure 2, (b). This distribution is also

¹⁵This statement holds also if voters anticipate valence shocks, because the utility difference between platforms, which depends on said projection only, is still a sufficient statistic for voting. It also holds, when instead of valence shocks, voting-cost shocks are present, provided they induce what is called abstention due to indifference (e.g. Ledyard, 1984) and not abstention due to alienation (Smithies, 1941). In the former case, the utility difference between parties is again a sufficient statistic for behavior, because agents vote if the utility difference exceeds the voting cost. Under abstention due to alienation, voters care not only about the relative but also about the absolute utility from parties, so they would be motivated to learn how far they are from the closer party platform.

Bayes-consistent due to the spherical prior. It induces the same projection of the posterior mean on the platform difference (or equivalently, on the diagonal line) and hence has the same instrumental value.¹⁶ And the Kullback-Leibler divergence is invariant under coordinate transformations, so reflections preserve the cost of information. For the second step, because the voter’s information cost is posterior separable, she is indifferent to mixing between equivalent distributions over posteriors and hence to acquiring the mixed distribution in Figure 2, (c), instead. For the third and last step, the voter prefers to acquire the distribution in Figure 2, (d), which presents a mean-preserving contraction of the mixed distribution and hence a garbling of the information. Thus, this distribution is cheaper to acquire while having the same instrumental value, because it has the same projection of posterior means on the platform difference. By symmetry of the mixed distribution constructed by the second step, the mean-preserving contraction in the third step results in a distribution supported on the line through the prior mean with direction $x_b - x_a$.

Theorem 1 is related to but distinct from two other results in the rational-inattention literature that can explain a reduction of dimensionality: learning about the partition of payoff-equivalent states only and the so-called water-filling algorithm.

Under rational inattention with an entropy-based cost, agents learn only about the partition of payoff-equivalent states (Sims, 2003; Caplin, Dean, and Leahy, 2022), which implies agents neglect payoff-irrelevant dimensions. This result does not necessarily hold when the information cost depends on the distance between states, as in recent contributions to the literature (Hébert and Woodford, 2021; Pomatto, Strack, and Tamuz, 2023). A concern about a result based on this mechanism is that it may require that voters are able to differentiate well between arbitrarily close states. By contrast, our proof builds on reflection-invariance of the information cost and holds for some plausible distance-based information costs as well.¹⁷ Furthermore, we show the induced distribution over posterior means is supported on a certain line, which makes predicted survey behavior indistinguishable from that under a one-dimensional policy space (Proposition 2), as used in much of formal political economy.

Theorem 1 is also reminiscent of the water-filling algorithm, which applies in linear-quadratic Gaussian tracking problems, that is, decision problems where agents choose a continuous action $x \in \mathbb{R}^n$ to track the state $\theta \in \mathbb{R}^n$ under a quadratic loss, $u(x, \theta) = -(x - \theta)^\top A(x - \theta)$, and a normal prior (Kőszegi and Matějka, 2020). According to the water-filling algorithm, attention is allocated to a subset of dimensions according to a particular order of priority, which is determined by how payoff-relevant these dimensions are. Further, the agent pays attention to more dimensions when the attention cost is lowered. By contrast, in our case, agents learn about at most one dimension, regardless of the information-cost parameter, because this dimension is sufficient for

¹⁶For general Σ and A , we construct a reflection that preserves both the elliptical prior and the payoff-relevant projection on the platform difference $x_b - x_a$ with respect to A .

¹⁷We show in Corollary 6 in Appendix D.3 that Theorem 1 holds for appropriate versions of the posterior-variance, neighborhood-based (Hébert and Woodford, 2021), and log-likelihood-ratio (Pomatto, Strack, and Tamuz, 2023) costs, which are all distance-based. The important condition is that the information cost is preserved under reflections that preserve the prior. Thus, our result does not allow for comparative statics under changing the information-cost distance and the prior separately.

decision-making purposes. The reason is that, in contrast to tracking problems, in our model, the agent can choose only from a discrete subset of the vector space.

Robustness and Extensions Theorem 1 is robust to several generalizations. The proof works for any distribution of valence ν , under correlated ideal points across voters, and under a heterogeneous information cost parameter κ in the population if κ is independent of ideal points (otherwise, voters could infer something about their ideal points from observing κ). Although our stark result relies on ex-ante homogeneity of voters and the existence of only two parties, appropriate extensions hold when we drop these assumptions.

First, the analysis can be extended to heterogeneous priors. One way to model heterogeneous priors is to suppose voters start with some homogeneous prior but obtain exogenous signals before learning, which creates heterogeneous interim beliefs. In particular, suppose voters have a normal prior and obtain normal signals $S_i = \theta_i + \varepsilon_i$, where the noise term ε_i has a common normal distribution and is independent of θ_i . Then, the interim beliefs after observing the exogenous signals are normal, that is, elliptical, with a covariance matrix that is common to all voters. Thus, Theorem 1 applies and all voters learn in the same direction. This implies learning does not change the marginal of the distribution of revealed ideal points on the hyperplane *orthogonal* to the direction $\Sigma A(x_b - x_a)$. By contrast, the marginal on the line increases in the mean-preserving spread order through learning. Thus, learning simply “stretches out” the distribution of revealed ideal points in the direction $\Sigma A(x_b - x_a)$ and does not increase the mean-squared error of predicting voter ideal points through their projection on said line.

Second, the result can be generalized to more than two parties. First, under plurality rule, in a Duvergerian equilibrium where voters decide between the two front-runners, our mechanism would still apply.¹⁸ Second, under electoral rules featuring proportional representation instead of plurality rule, expressive voting may remain a good assumption. Maintaining expressive voting and assuming $k > 2$ parties, we show by an analogous reflection argument that voters’ ideal points lie on an at most $(k - 1)$ -dimensional hyperplane. If k platforms were to lie on a line, however, voter ideology would still be one-dimensional.¹⁹ If the platforms lie in general position, the model predicts that the dimensionality of voter ideology is increasing in the number of parties. Unfortunately, we are not aware of any systematic evidence regarding this prediction.

Corollary 1. *If voters face k party platforms, the distribution of revealed ideal points ρ has support inside a subspace of dimension at most $k - 1$.*

This result aligns with evidence from Western European countries, which have more than two parties and where a two-dimensional political space is needed to capture party and voter positions (Kriesi, Grande, Lachat, Dolezal, Bornschier, and Frey, 2006; Kriesi, Grande, Lachat, Dolezal,

¹⁸Although we assume expressive voting, strategic voting does not undermine our result under private values as long as voters perceive a positive probability of being pivotal.

¹⁹The dimensionality of voter ideology may also be strictly smaller than the number of parties minus 1 if voters neglect some party (e.g., because it is too far away from their prior mean) in their learning problem. In the terminology of Caplin, Dean, and Leahy (2019), this happens if a party is not in the voter’s consideration set. It follows that the dimensionality of voter ideology is less than or equal to the size of their consideration set minus 1.

Bornschier, and Frey, 2008; Bornschier, 2010a; Bornschier, 2010b; however, see Van Der Brug and Van Spanje, 2009).

3.2 Polarized Ideology

Our second main result shows that without uncertainty about valence, the optimal signal structure induces a binary distribution of revealed ideology. This result holds even if the true distribution of ideal points is continuous and unimodal. With uncertainty about valence, this result does not necessarily hold, but we show for “small” valence shocks, the distribution of revealed ideal points is “almost” binary.

Theorem 2 (Polarization). *Absent valence shocks, voters’ revealed ideal points are supported on at most two points. If the distribution of the valence shock converges in mean to zero, any selection of optimal distributions over revealed ideal points converges weakly to a binary distribution.*

Without valence shocks, voters only want to learn what party they are closer to. In other words, the voter faces a binary decision problem after learning. In rational-inattention problems with k actions, an optimal signal structure that induces at most k distinct posteriors is known to exist.²⁰ The reason is that if the agent acquired more posteriors—and thus signals—than actions, they could garble the signal structure based on the action recommendation. This garbling would maintain the instrumental value of information while saving on the information cost, because the garbling leads to a Blackwell-dominated signal structure.

To our knowledge, the rational-inattention literature has not emphasized the implication of this result for polarization. Under rational inattention, the necessity to take an action makes agents learn about their preferences in a way that divides them into discrete groups, one for each action—or, in our case, one for each party. In fact, this mechanism holds under flexible information acquisition for any strictly Blackwell-monotone information cost, that is, for any cost that makes a strict garbling of the signal structure strictly cheaper.

When valence shocks realize after learning, they effectively enlarge the choice set and break the mechanism for binary learning. The choice set is larger because the voter can now decide for each realization of valence ν who to vote for. Or, equivalently, voters now care about learning *how much* they prefer one party to the other. Such learning informs them for what size of the valence shock they should start voting for party b . This results in a continuous rational-inattention problem, which generally do not have closed-form solutions (Jung, Kim, Matějka, and Sims, 2019).

However, Theorem 2 shows a continuity result for valence shocks close to degenerate. If the valence shock converges in mean to zero, the distribution over revealed ideal points converges to a binary distribution. This result implies that for any two open neighborhoods of the two points of the binary distribution, as valence converges to zero, the mass of these two neighborhoods converges to 1. That is, for small-enough valence, almost all revealed ideal points will be very close to one of the two points, so we can talk essentially of a bimodal distribution.

²⁰For the Shannon-entropy cost, the result has been observed by Sims (2003). In Bayesian persuasion, it has been observed by Kamenica and Gentzkow (2011).

To show the continuity result, we prove a more general continuity result for information design problems in Appendix D (Proposition 7). Proposition 7 applies to any information design problem with state space $\mathbb{R}^n$ and an upper semicontinuous value function that is bounded from above. It establishes that the solution is upper hemicontinuous in the topology of weak convergence under uniform convergence of the value function, which may be useful beyond our application. The proof of Proposition 7 is complicated by the fact that, unlike existing result (Caplin, Dean, and Leahy, 2022; Dworzak and Kolotilin, 2023), we do not restrict ourselves to a finite or compact state space. Moreover, we cannot assume a continuous value function, because with an infinite state space, the Kullback-Leibler divergence is only lower semicontinuous rather than continuous. We show that using a generalization of Berge’s maximum theorem due to Tian and Zhou (1992), we can nonetheless obtain our result. Our result shows a sense in which it is not true that “similar decision problems may lead to sharply different [behavior]” (Jung, Kim, Matějka, and Sims, 2019), which is reassuring for the theory of rational inattention.

Comparing our result with the well-studied tracking problems in the rational-inattention literature, introduced above, is again instructive. Under quadratic loss and normal prior, the agent is known to optimally acquire a normal signal, resulting in a normal distribution over posterior means that, of course, cannot be bimodal. Although the presence of continuous valence shocks makes voters’ choice set effectively continuous, the utility is not quadratic in the distance of action and state, which allows for a *bimodal* distribution over posterior means. Relatedly, Jung, Kim, Matějka, and Sims (2019) show in tracking problems, when the utility depends on the distance between the action and the state but not in a quadratic way, agents will often choose from a discrete set of actions only. In our case, the utility is not a function of the distance between action and state, so their result does not apply. Instead, the bimodality is driven by the existence of two underlying options, as explained above.

Two related results study belief polarization over a common state. Nimark and Sundaresan (2019) show that the beliefs of a population of rationally inattentive agents can become polarized over time, as agents information acquisition is path-dependent. Eguia and Hu (2022) show beliefs can become polarized if agents are boundedly rational in the sense of a finite memory and have heterogeneous preferences. Our result that revealed voter ideal points are binary without valence also holds if ideal points are common. We expect that the continuity result also generalizes, but this part requires additional work. By comparison to above papers, we show a polarized distribution can result without ex-ante heterogeneity and dynamics or bounded rationality.

3.3 Comparative Statics

Broadly speaking, polarization of a distribution is understood as capturing how bimodal and how spread out the distribution is (Esteban and Ray, 2012). We have shown above flexible information acquisition predicts bimodal ideology when valence shocks are small. Here, we show how a smaller cost of information or more distant party platforms can polarize voters, in the sense of leading to a more spread out distribution of revealed ideal points. We use this comparative statics result for

our third main result, Theorem 3.

For the comparative statics result, we assume a normal prior and restrict voter learning to *normal* signals while maintaining that the information cost is proportional to mutual information. Formally, we define a normal signal as a random vector S such that (S, θ) is jointly normal. We make this simplification because comparative statics under flexible information acquisition are notoriously difficult due to the high dimensionality of the signal choice.²¹ By contrast, under the restriction to normal signals, and because voters learn only in a one-dimensional way, their candidate signals are completely Blackwell-ordered, which we exploit for the proof.

For this comparative statics result, we also assume the party platforms are equally distant from the voter’s prior mean under the distance relevant to voter preferences, $x_a^\top A x_a = x_b^\top A x_b$, an assumption we revisit in section 4.1. We formalize this by supposing that party platforms (x_a, x_b) are a scaled version of (x, y) , $(x_a, x_b) = \alpha(x, y)$, with $x^\top A x = y^\top A y$. We vary the scalar α , which we call the degree of platform polarization.

Proposition 1 (Comparative Statics). *Restrict the prior μ and feasible voter signals to be normal and let $(x_a, x_b) = (\alpha x, \alpha y)$ with $\alpha \in \mathbb{R}_{\geq 0}$. The variance of the optimal distribution of revealed ideal points strictly increases in the strong set order when the information cost parameter κ decreases and when the degree of platform polarization α increases.*

Because the optimal signal structure may not be unique, the comparative statics result is expressed in terms of the standard strong set order. The intuition is as follows.

First, smaller κ or larger α encourage voters to acquire more information. As information becomes cheaper, voters learn more by supermodularity of their objective in the parameter κ and the cost of information $c(\tau)$, using the fact that the candidate signal structures are completely Blackwell ordered. If party platforms were very close to each other, it would not matter much for voters who to vote for, so they would learn little about their ideal points. As party platforms are more polarized, voters face larger stakes in the election and acquire more informative signals.

Second, more information leads to a distribution of voter ideal points with higher variance. While one might expect that more information leads to more agreement, here it leads to more disagreement simply because voters learn about their idiosyncratic ideal points. A more informative signal leads to a mean-preserving spread of the distribution of posterior means. Because voters learn about their *independent* ideal points, this translates to a mean-preserving spread of revealed ideology. We show, after the proof of Proposition 1 in Appendix D, that this conclusion is robust to some correlation between ideal points through a common component. The result is robust as long as the variance of the common component is smaller than the variance of the idiosyncratic component.

²¹Yoder (2022), who provides a comparative statics result under a small state space, notes that the value and cost of information need not be quasipermodular in τ , so one cannot apply the comparative statics by Milgrom and Shannon (1994). See also the discussion in Curello and Sinander (2024) on costly information acquisition, which shows that even under posterior-mean separable information costs, increasing comparative statics hold only under very special conditions.

The prediction that more information leads to greater polarization is consistent with evidence. Palfrey and Poole (1987) develop an index of voter information and find more informed voters tend to be more extreme. Abramowitz and Saunders (2008) and Abramowitz (2010) find that more educated and engaged voters are more ideologically extreme. Lauderdale (2013) provides causal evidence that increasing information leads to ideological polarization. We discuss the evidence on the comparative statics regarding platform polarization in section 3.4.3.

A sizeable literature studies how information can lead beliefs about a *common* state to diverge. The beliefs of agents with heterogeneous priors can diverge when observing a common signal, due to ambiguity aversion (Baliga, Hanany, and Klibanoff, 2013) or uncertainty about the signal structure (Acemoglu, Chernozhukov, and Yildiz, 2016). Novák, Matveenko, and Ravaioli (2024) also studies rationally inattentive agents, but with a common prior and heterogeneous preferences for the status quo in a binary decision problem. They show beliefs may diverge in expectation, conditional on the true state of the world, as agents acquire different signal structures. In contrast to these papers, our agents learn about *idiosyncratic states*, namely, their independent ideal points. However, as mentioned above, our monotone comparative statics would also hold in the presence of a common component of ideal points, provided the variance of the common component is smaller than the variance of the idiosyncratic component. Our focus on idiosyncratic ideal points is motivated by our application. While above papers aim at explaining persistent disagreement about facts, we focus on political positions, which are naturally heterogeneous due to conflicting interests. We therefore take seriously that voters need to learn about idiosyncratic factors affecting their political positions. This provides a simple and natural explanation of how information leads to increasing spread of voter ideal points.

3.4 Evidence

We relate our results to the existing evidence on voter ideology.

3.4.1 Issue Alignment

Recent evidence shows that the ideology of US voters is approximately one-dimensional (Jessee, 2009; Jessee, 2012; Tausanovitch and Warshaw, 2012; Shor and Rogowski, 2018; Fowler, Hill, Lewis, Tausanovitch, Vavreck, and Warshaw, 2022; Hare, Highton, and Jones, 2023). These studies use voter surveys, such as the American National Election Studies, to estimate voter ideal points, similar to the ideal point estimation of legislators from roll-call data (Poole and Rosenthal, 1985). Specifically, these studies estimate a one-dimensional spatial model with quadratic utility to predict the binary responses $y_{ij} \in \{0, 1\}$ of each individual i to each question j (e.g. should the minimum wage be raised).²² According to the model, the likelihood is $\Pr(y_{ij} = 1) = \Phi(u(x_{j1}, \theta_i) - u(x_{j2}, \theta_i))$,

²²To be even more precise, these models estimate one-dimensional item-response theory models, which are known to be equivalent to one-dimensional spatial models with quadratic utility (e.g. Ladha, 1991). The only exception is Hare, Highton, and Jones (2023), who use a different methodology but also conclude that voter ideology is approximately one-dimensional.

where Φ is the logistic or normal cumulative distribution function and utility is quadratic, $u(x, \theta) = -(x - \theta)^2$. The to-be-estimated parameters are the ideal points $\theta_i \in \mathbb{R}$ of each individual i and the positions $x_{j1}, x_{j2} \in \mathbb{R}$ of the policies corresponding to two responses of each question j (e.g. a minimum wage raise and the status quo). This is a standard logit or probit discrete choice model, where a voter responds more likely with the policy closer to their ideal point. These studies find that such a one-dimensional model explains voter responses well (typically about 80% of binary responses) and that adding more dimensions only marginally increases the explanatory power of the model. They conclude that ideology is well described by a one-dimensional ideological spectrum.²³

Upon closer examination, it is not clear whether the prediction of Theorem 1 aligns with the evidence that survey responses are well explained by a one-dimensional spatial model. First, would it not be sufficient for voter ideal points to be on some one-dimensional *curve* for a one-dimensional spatial model to explain voter's survey responses? In that case, Theorem 1 would be proving too much. Second, does equivalence to a one-dimensional spatial model require not just that the ideal points but also the *policies* are in a one-dimensional space? Theorem 1 predicts one-dimensional ideal points within a multidimensional policy space, while in one-dimensional spatial models *both* the ideal points as well as the policies live in a one-dimensional space. If the answer to the second question is affirmative, then Theorem 1 would be proving too little to explain the evidence.

In the following, we show neither is the case and Theorem 1 proves the property of voter ideology identified by the evidence, namely the property that ensures that voters' survey responses can be explained by a one-dimensional spatial model.

First, we need additional definitions. A *multidimensional spatial model* with quadratic utility is defined identically to the one-dimensional spatial model described above, except for the parameters $\{\theta_i, x_{j1}, x_{j2}\}$ being elements of $\mathbb{R}^n$ and $u(x, \theta) = -(x - \theta)^\top (x - \theta)$ being the multidimensional analogue of quadratic utility.²⁴ It turns out that the property of a multidimensional spatial model identified by the evidence is that respondents' *ideal points are on a line when projected onto the space spanned by the survey questions*. Formally, this property states that there exist $\lambda_i \in \mathbb{R}$, $\Delta\theta \in \mathbb{R}^n$ and $\theta_i^\perp \in \mathbb{R}^n$, such that

$$\begin{aligned}\forall i: \theta_i &= \theta_1 + \lambda_i \Delta\theta + \theta_i^\perp, \\ \forall i, j: (x_{j1} - x_{j2})^\top \theta_i^\perp &= 0.\end{aligned}$$

²³This finding stands in contrast to the older literature, starting with Converse (1964), which studies correlation between voter responses to different policy questions instead of estimating ideal points. These papers typically find low correlation and conclude that ideology not well represented by a one-dimensional spectrum, or that there is little *constraint* on voter ideology in the terminology of Converse. Later research has found that this conclusion is partly driven by response mistakes such as arising from inattentiveness of respondents, which reduce correlation (Ansolabehere, Rodden, and Snyder, 2008). Further, the literature seems to have overlooked another reason for why the correlation between responses is a poor measure of to what extent voter ideology is one-dimensional. Even if voters respond to questions according to a one-dimensional spatial model, they may not consistently give left or right responses but respond with whichever response option is closer to their ideal point. In the one-dimensional spatial model, the response to question j depends on whether the voter's ideal point θ is below or above the question midpoint $\frac{1}{2}(x_{j1} + x_{j2})$. If different questions have different midpoints, the voter would be expected to choose the left or right response depending on the question.

²⁴The result remains the same if we assume a general quadratic form $u(x, \theta) = (x - \theta)^\top A(x - \theta)$ because we can switch to an orthonormal basis of A .

That is, each ideal point θ_i is on the line through θ_1 with direction $\Delta\theta$, modulo a component $\theta_i^\perp$ that is orthogonal to the policy-differences $x_{j1} - x_{j2}$ for each question j . We say that a multidimensional spatial model with (n -dimensional) parameters $\{\theta_i, x_{j1}, x_{j2}\}$ is *observationally equivalent* to a one-dimensional spatial with (one-dimensional) parameters $\{\hat{\theta}_i, \hat{x}_{j1}, \hat{x}_{j2}\}$ if they predict the same likelihood $\Pr(y_{ij} = 1)$ over survey responses for all i and j .

Proposition 2. *Under quadratic utility, a multidimensional spatial model is observationally equivalent to some one-dimensional spatial model if and only if the multidimensional ideal points are on a line when projected onto the space spanned by the survey questions.*

Proposition 2 shows ideal points being on a line (when projected onto the space spanned by the survey questions) is the property of ideal points that makes survey behavior explainable by a one-dimensional spatial model. The parenthesized caveat holds because, naturally, survey responses are not affected by policy dimensions that are orthogonal to all survey questions. Because voter surveys try to cover most relevant policy dimensions, we take this caveat to be of limited importance. Since the above-mentioned papers show voters’ survey responses are well-explained by a one-dimensional spatial model (and not much better by higher-dimensional models), we conclude they confirm the prediction of Theorem 1.

The high-level intuition for Proposition 2 is as follows. For the “only if”-direction, suppose voter ideal points were not on a line but, say, on a U-shaped curve. Then, the extreme voters on both sides of the U may prefer some policy to a policy at the bottom of the U, that is preferred by the centrist voters. Such non-monotonic behavior is ruled out by one-dimensional ideological spectrum. For the “if”-direction, for any survey question, one can find suitable projections of its two policies onto the voter line that do not change predicted behavior and make the model one-dimensional.

3.4.2 Polarized Ideology

Whether voters are ideologically polarized, that is, have a bimodal distribution, has lead to an academic debate between Abramowitz and Saunders (e.g. Abramowitz and Saunders, 2008) on the affirmative side, and Fiorina, Abrams, Pope, and Levendusky (e.g., Fiorina and Abrams, 2008) on the other (see Lelkes, 2016 for a critical overview of the debate). Unfortunately, neither side of the debate estimates voter ideal points but only uses “raw” survey evidence, so it is not clear how to interpret their findings. For example, much of the evidence against ideological polarization stems from evidence on ideological self-placements on 7-point scales (e.g. Fiorina and Abrams, 2008). The validity of ideological self-placements has been criticized for several reasons but the literature has, to our knowledge, overlooked a more fundamental problem. A 7-point scale is a categorical, *ordinal* scale. To assess the bimodality of the distribution of voter ideology, a *cardinal* scale is needed. The reason is that one can always monotonically transform the scale to make a distribution bimodal or unimodal. An ideological 7-point scale is only meaningful if one assumes the 7 categories correspond to intervals of the same size on the appropriate cardinal scale of ideology, for which the authors provide no evidence. This underlines the importance of estimating ideal

points from survey responses to multiple questions, as this results on ideal points on a meaningful cardinal scale. On the other hand, the evidence *for* polarization relies on measures of individual-level correlation between left and right responses to different policy questions (Abramowitz and Saunders, 2008). However, without estimation of ideal points it is not clear whether their findings relate to issue alignment or polarization of ideology.

Several newer papers do, however, estimate voter ideal points from survey responses. Bafumi and Herron (2010) find a bimodal distribution of voter ideal points, while most papers find a unimodal distribution (e.g., Jessee, 2012; Hill and Tausanovitch, 2015; Dun and Jessee, 2020). However, there are reasons to believe that current estimation algorithms underestimate ideological polarization of voters. Survey respondents who do not pay much attention to the survey are artificially placed in the middle of the distribution, because this best explains their random responses (McCarty, 2019, 204). Indeed, when Fowler, Hill, Lewis, Tausanovitch, Vavreck, and Warshaw (2022) screen for inattentive respondents (and for respondents that are not well-represented by a one-dimensional ideal point), they find a more bimodal distributions of ideal points in most survey years. Moreover, Abramowitz (2010) finds the distribution of actual or engaged voters is more polarized. We conclude that the matter is not settled yet.

3.4.3 Party Influences on Voter Ideology

Political scientists have long argued that mass opinion is heavily influenced by the elite political discourse (Zaller, 1992; Lenz, 2012), yet the underlying mechanisms remain debated (Leeper and Slothuus, 2014). Our model provides a mechanism through which both issue alignment and polarization of voters is affected by parties. Perhaps surprisingly, this mechanism is consistent with voter rationality.

Specific to issue alignment, Malka, Lelkes, and Soto (2019) write “political scientists generally agree that [issue alignment] among politically attentive citizens results from such citizens following elite political cues.” This idea is also motivated by findings such as that the meaning of left and right changes over time and space, in congruence with party positions (Inglehart and Klingemann, 1976). For example, whether protectionism is associated with the left or right in the US has evolved over time (McCarty, 2011). Theorem 1 predicts the orientation of the ideological spectrum, that is, $\Sigma A(x_b - x_a)$, is determined by party platforms, x_a and x_b . As argued in the introduction, this orientation determines what issues go together.

Implicit in the understanding by Malka, Lelkes, and Soto (2019) is that the issue alignment among voters is consistent with relative party platforms. That is, for example, if one party is more left on economics and more liberal on social issues than the other party, then voters who are more left are also more liberal. More precisely, we say *issue alignment is consistent with party platforms* if the sign of the k -th component of relative party platforms, $x_b - x_a$, equals the sign of the k -th component of the orientation of the ideological spectrum, $\Sigma A(x_b - x_a)$, for all $k = 1, \dots, n$. Then, for any two dimensions the alignment of relative party platforms conforms to the issue alignment of voters. While this is not a necessary prediction of our model, it holds in important special cases.

If $x_b - x_a$ is an eigenvector of ΣA , then the ideological spectrum is exactly parallel to the platform difference $x_b - x_a$, so issue alignment is consistent with party platforms. In section 6, we give a microfoundation for $x_b - x_a$ being an eigenvector of ΣA if party objectives are driven by valence competition. This alignment may also occur in a richer model in which party ideology arises from the ideology of voters who are party members, resulting in party platforms that are on the line of voter ideal points. If the covariance matrix Σ of true ideal points and the matrix A associated with the policy utility are both diagonal, the issue alignment is also consistent with party platforms.²⁵ Diagonality holds if voter positions on different policy issues are uncorrelated and there are no preference interdependencies between issues. Broadly speaking, as long as such correlations and interdependencies are not strongly enough misaligned with relative party positions, issue alignment should be expected to be consistent with party platforms.

Regarding polarization, Proposition 1 shows how platform polarization can lead to polarization of voters. This is not because voters blindly follow party positions but instead as a consequence of rational learning. Again, the political environment can affect revealed ideology even when true ideology remains unchanged. This is consistent with the finding of Bischof and Wagner (2019) that voters ideal points diverge immediately after new radical parties enter parliament.

4 Electoral Competition

We are interested in welfare properties of the equilibrium platforms, from the viewpoint of voters' policy utility. For voters, it is crucial how much parties polarize their platforms, moving away from the policy that maximizes voter's aggregate policy utility. Before we turn to this question in section 4.1, we highlight some important forces at place in electoral competition, holding voter preferences fixed.

Recall that parties choose their platforms, x_a and x_b , in a Nash equilibrium of the electoral-competition game given the distribution ρ of voters' revealed ideal points. It is typically hard to obtain characterizations of equilibrium platforms when parties are motivated both by vote share and ideology. However, our party objective, which is linear in vote share and ideological utility, allows such a characterization. The following lemma shows (1) party platforms are a weighted mean of voter and party ideal points, and (2) the weight on a voter is decreasing in the "extremeness" of the voter. We use these properties subsequently and relate them to platform polarization.

Lemma 1. *In any equilibrium, party platforms are a weighted average of voter and party ideal points,*

$$x_a = \frac{m \int w(\theta) \theta d\rho(\theta) + x_a^*}{m \int w(\theta) d\rho(\theta) + 1}, \quad (7)$$

$$x_b = \frac{m \int w(\theta) \theta d\rho(\theta) + x_b^*}{m \int w(\theta) d\rho(\theta) + 1}. \quad (8)$$

²⁵If $\Sigma = \text{diag}(\Sigma_{11}, \dots, \Sigma_{nn})$ and $A = \text{diag}(A_{11}, \dots, A_{nn})$, with $\Sigma_{11}, \dots, \Sigma_{nn}, A_{11}, \dots, A_{nn} > 0$ by positive definiteness, the k -th component of $\Sigma A(x_b - x_a)$ is simply $\Sigma_{kk} A_{kk}(x_{b,k} - x_{a,k})$, which has the same sign as $x_{b,k} - x_{a,k}$ for $k = 1, \dots, n$.

where

$$w(\theta) = f_\nu(u(x_a, \theta) - u(x_b, \theta)).$$

The result generalizes the mean-voter theorem by Hinich (1977), which assumes purely office-motivated candidates. The mean-voter theorem states that under quadratic voter utility and probabilistic voting, party platforms converge at the mean of voter ideal points. Because parties are office- and ideologically-motivated in our model, their platforms are affected by voter ideal points as well as the party’s own ideal point in an intuitive way.

In particular, by symmetry of f_ν , the weight $w(\theta)$ on a voter with ideal point θ depends only the size of the utility difference between party platforms, $|u(x_a, \theta) - u(x_b, \theta)|$. Voters that have a larger utility difference can be seen as more “extreme” or “ideologically entrenched” relative to the party platforms. By strict quasi-concavity of f_ν , a voter with a larger utility difference has a smaller weight $w(\theta)$. Intuitively, more extreme voters are less sensitive to platform changes (the probability that they change their vote due to a small platform change is small), so they have less influence on equilibrium platforms. While this observation is not new (Persson and Tabellini, 2002, 57), most models of probabilistic voting rule this effect out by focusing on a uniform distribution of valence shocks to improve tractability. We use this observation later to show voter polarization amplifies platform polarization: If voters are more extreme on average, parties moderate less and choose policies closer to their own ideal points. This mechanism is consistent with the finding by McCarty, Rodden, Shor, Tausanovitch, and Warshaw (2019) that more ideologically heterogeneous districts have more extreme legislators.²⁶

Lemma 1 only speaks to *necessary* conditions of equilibrium platforms, as derived from first-order conditions. Hence, additional work is necessary to show that equilibrium candidates that satisfy the first-order conditions constitute actual equilibria. For example, they constitute equilibria if the party objectives are quasi-concave, in which case the first-order conditions are sufficient for optimality. Along these lines, in Appendix D.7, we give a condition that ensures that our equilibrium candidates in the context of Theorem 3 are equilibria. We also show this condition is satisfied when the weight m on vote share is small enough or when the valence shock ν is large enough, echoing observations by Lindbeck and Weibull (1987) and Enelow and Hinich (1989).

4.1 Platform Polarization

To illustrate the mechanisms underlying platform polarization, we simplify to a symmetric setup. A symmetric setup is obtained when the party ideal points x_a^* and x_b^* are equally far from the origin according to the distance relevant to voter preferences, $x_a^{*T} A x_a^* = x_b^{*T} A x_b^*$. We show below under this assumption all equilibria are symmetric in the following sense. We say (ρ, x_a, x_b) is a *symmetric* equilibrium if

$$(x_a, x_b) = (\alpha x_a^*, \alpha x_b^*)$$

²⁶McCarty, Rodden, Shor, Tausanovitch, and Warshaw (2019) interpret this finding through the Calvert-Wittman model, in which greater uncertainty about the location of the median voter leads to greater platform polarization. They theoretically connection voter polarization to uncertainty about the median voter through the informativeness of a public poll. Our model provides a more direct mechanism.

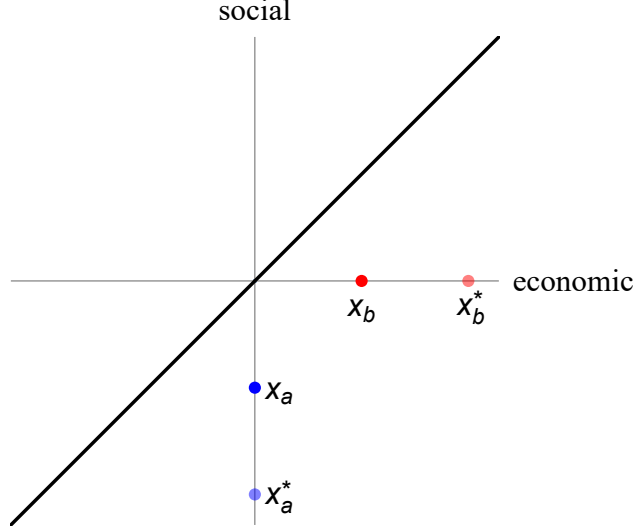


Figure 3: Symmetric equilibrium with party platforms (x_a, x_b) and party ideal points (x_a^*, x_b^*)

with $\alpha \in (0, 1)$, which we call the *degree of platform polarization*. Theorem 3 below identifies comparative statics for the degree of platform polarization α as the information cost parameter κ changes. Figure 3 visualizes an example of a symmetric equilibrium. Note that a higher degree of platform polarization not only increases the distance between party platforms but also makes party platforms move further away from the ideological spectrum of voter ideal points.

Bringing together endogenous voter ideology and endogenous party platforms, we show the following result. We restrict again to normal distributions to make use of the comparative statics result Proposition 1. Because voter and platform polarization are mutually reinforcing, there may be multiple equilibria, which we order by their degrees of platform polarization α . Therefore, as usual, our comparative statics are expressed in terms of the smallest and largest equilibrium.

Theorem 3. *Restrict the prior μ and feasible voter signals to be normal. There exists an equilibrium and every equilibrium is symmetric. Cheaper information increases polarization: The smallest and largest equilibrium degree of platform polarization α weakly increase as κ decreases.*

Theorem 3 combines our earlier results on voter ideology and on party platforms. To prove the theorem, we show voter polarization and platform polarization are mutually reinforcing: if voters are more extreme, their voting is less sensitive to party platforms, allowing parties to polarize more (Lemma 1). If platforms are more polarized, then voters face larger stakes in the election, inducing them to learn better and become more extreme (Proposition 1). One can think of cheaper information to start this self-reinforcing process by allowing voters to learn at a lower cost (Proposition 1). On a formal level, we establish existence of pure-strategy equilibria and the comparative statics result, through monotonicity arguments similar to those in supermodular games (despite our game not being supermodular).

The theorem implies better availability of information makes the equilibrium policy worse for voters. Platform polarization hurts voters in our model because the utilitarian optimum for voters

is the policy coinciding with the mean ideal point, which is at the origin. A higher degree of platform polarization, α , implies that any implemented policy (x_a or x_b) moves further away from the origin, decreasing voter’s aggregate policy utility. While cheaper information allows voters to learn more accurately about their ideal points, this makes voters less responsive to party platforms, leading to greater platform polarization. Voters do not internalize this information externality of their learning strategy on party platforms because each voter is infinitesimal.

Theorem 3 underscores the different implications of learning about preferences versus learning about equilibrium actions of other agents. Matějka and Tabellini (2021) show more informed voters are *more* responsive to party platforms, when voters learn about party platforms knowing their ideal points. This would suggest decreasing platform polarization in equilibrium as information becomes more accessible. By contrast, in our model, better informed voters are more extreme and therefore *less* responsive to party platforms. Furthermore, we show in section 5 this difference is not due to our timing assumption. In a symmetric equilibrium, the vote share is less responsive to the platform choice under cheaper information, also if parties publicly commit to their platforms *before* voters learn about their preferences.

The theorem demonstrates one mechanism that may have contributed to increasing party polarization in the US, as observed in the past decades (McCarty, Poole, and Rosenthal, 2016). Information can become cheaper due to advances in information technology, such as the internet. Theorem 3 shows better availability of information can lead to more platform polarization. The underlying mechanism operates through increasing polarization of voters. While the empirical evidence on increasing polarization of US voters is somewhat mixed, it suggests that voter polarization may have increased more recently. Hill and Tausanovitch (2015) find that the variance of US voter ideal points is generally stable from 1956 to 2012, but their point estimates for variance increase after the year 2000.²⁷ The Pew Research Center (2014) find similar spread of voter position in 1994 and 2004 but a significant increase in 2014. Thus, one may take Theorem 3 to suggest a contributor to platform polarization in the post-2000 era.

4.2 Aggregate Uncertainty

Until now, we have assumed that there is no aggregate uncertainty about voter preferences. Therefore, the optimal policy was always at the mean of voter ideal points, that is, the origin. In general, however, the optimal policy may depend on such aggregate uncertainty. This opens up the new question whether under endogenous voter learning, elections aggregate preferences, in the sense of making policy responsive to aggregate uncertainty. We show a novel failure of information aggregation: because voter learning is one-dimensional, policy responds to only one dimension of aggregate uncertainty.

We model aggregate uncertainty about voters’ ideal points through an aggregate state ω , which enters voters ideal points as a common component. The ideal point of voter i is $\theta_i = \omega + \delta_i$, where the

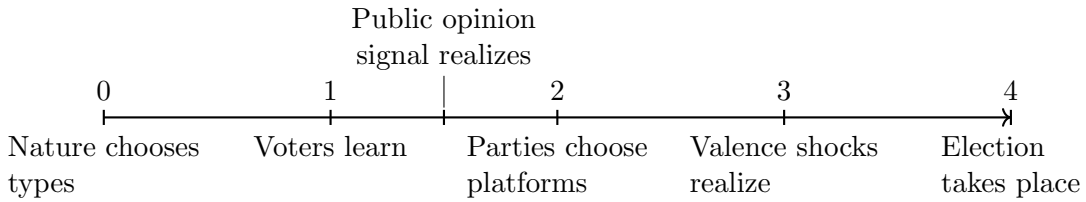
²⁷Moreover, in subsequent unpublished work, they find more evidence of increasing polarization in the post-2012 data, but the conclusion depends on the measure and statistical model (personal communication).

idiosyncratic components $\{\delta_i\}$ are identically distributed and $(\omega, \{\delta_i\})$ are jointly independent. For our result in this section, we do not need to impose elliptical distributions for ω and δ_i . However, to simplify the proof of equilibrium existence and because it comes with little loss of economic substance, we assume the support Ω of ω and the support D of δ_i are finite. We allow voters i to acquire any signal structure about (ω, δ_i) and maintain the assumption that all voters acquire the same signal structure. Formally, each voter $i \in [0, 1]$ chooses a signal structure, that is, a stochastic kernel $\sigma_i : \Omega \times D \rightarrow \Delta(\mathcal{S})$ that maps each state (ω, δ_i) into a distribution over signals of a sufficiently rich signal space $\mathcal{S}$.²⁸ We understand the common component ω as a way to model aggregate uncertainty about voter preferences and assume that it does not affect party ideal points.

For parties to be able to respond to realized voter preferences, they must obtain information about voter preferences. If parties had *private* information about the realized distribution of revealed voter ideology, then platforms could convey information about the common component ω to voters. This would introduce a signalling motive into electoral competition (see Martinelli, 2001). However, since this signaling motive is not the focus of our analysis, we assume instead that both parties and voters observe a *public* signal $s \in S$ about voter preferences. This public signal could represent a poll or, more generally, any channel through which information about public opinion is disseminated.

Formally, the public signal is a stochastic kernel $\sigma_p : \Delta(\mathcal{S}) \rightarrow \Delta(S)$ that maps the realized distribution over voters' private signals into a distribution over public signals. As the public signal depends only on the distribution of voters' private signals, the realization of the public signal cannot be affected by a single infinitesimal voter. To show equilibrium existence, we assume the public signal space S is finite and the probability of any public signal $s \in S$, $\sigma_p(s|\pi)$, is continuous under weak convergence of the distribution $\pi \in \Delta(\mathcal{S})$ over private signals. This weak assumption is satisfied if the public signal contains some amount of noise.

Our extended game thus contains an additional stage between voter learning and platform choice where the public opinion signal realizes. As a result, parties can condition their platforms (x_a, x_b) on the realization of the public signal $s \in S$ and voters can condition their voting behavior both on the chosen platforms (x_a, x_b) and on the realization of the public signal s .²⁹



The main result of this section shows that while we introduce a channel through which party platforms can respond to the aggregate state ω , they respond to only one dimension of ω , namely the projection of ω on the ideological difference between parties.

²⁸While up till now, we have modelled signal structures as distributions over posteriors, to prove equilibrium existence in the context of this section, it is more useful to model signal structures as Blackwell experiments.

²⁹See the Appendix for a formal definition of strategies.

Theorem 4. *There is an equilibrium in which party platforms are affected by the aggregate state ω only through its A -projection on $x_b^* - x_a^*$, $(x_b^* - x_a^*)^\top A\omega$. That is, the distribution over equilibrium policy is the same under aggregate states ω and ω' if $(x_b^* - x_a^*)^\top A\omega = (x_b^* - x_a^*)^\top A\omega'$.*

The intuition is as follows. Parties only respond to components of the aggregate state ω that voters learn about. Consequently, if voters learn only about the A -projection of ω on the ideological difference $x_b^* - x_a^*$ between parties, then party platforms are only affected by this component. While Theorem 1 suggests that voters optimally learn only about one component of ω , the public opinion signal introduces two complications for showing this. First, at the time of learning, voters do not know yet what party platforms will be, giving them potentially an incentive to learn about multiple dimensions of their ideal points. However, independent of the public opinion signal, the platform difference $x_b - x_a$ is *parallel* to the ideological difference $x_b^* - x_a^*$ of parties, by Lemma 1. Therefore, voters have no incentive to learn about components of their ideal points orthogonal to $x_b^* - x_a^*$ for the purpose of voting. Second, the information about ω obtained through the public opinion signal could be complementary to private learning about orthogonal components of ω . However, we show that there exists an equilibrium in which no voter learns about such orthogonal components of ω , which rules out such complementarities. It is an open question whether there are equilibria in which party platforms respond to more than one component of ω .

Theorem 4 presents a severe inefficiency of preference aggregation due to endogenous voter learning. Because of independent idiosyncratic components δ_i , the average ideal point equals the common component ω . Because of quadratic preferences this makes ω the policy that maximizes unweighted aggregate voter welfare. However, voters learn only about the dimension of the common state ω along which parties disagree and, as a consequence, equilibrium policies respond only to this *one dimension* of ω , even if the policy space, and thus ω , is high-dimensional.

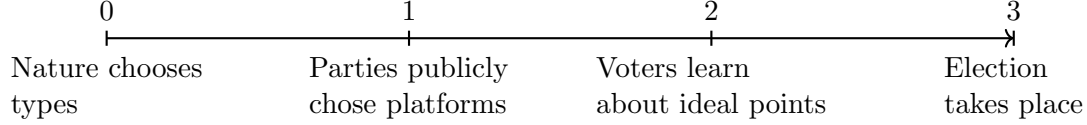
The result is particularly relevant because the failure of preference aggregation and resulting welfare loss might not show in the data and therefore go unnoticed. Judging on the basis of *revealed* ideology, it seems that parties do respond to voter preferences. However, the two-party system prevents voter learning about their preferences in more than one dimension, so there is a large scope for *unrevealed* voter ideology that policy does not respond to.

5 Alternative Timing: Parties as Agenda-Setters

5.1 Model

Game: This section considers an alternative timing where parties choose platforms $x_a, x_b \in \mathbb{R}^n$ before voters learn about their ideal points. Thus, parties act as agenda-setters, that is, through their chosen positions, parties affect what issues voters pay attention to. For example, if a party polarizes on migration, then voters will pay more attention to this issue because it becomes more relevant to their voting decision.

Voters: Voter preferences are as in the baseline model, except for the absence of a valence shock and we assume equal priorities, $A = I_n$. That is, the utility U_i of voter $i \in [0, 1]$ has two



components: Voter utility depends (i) on the implemented policy x and the voter's ideal point $\theta_i \in \mathbb{R}$ via the policy utility $u(x, \theta_i)$ and (ii) on the information cost $c(\tau_i)$ of signal structure $\tau_i \in \Delta(\Delta(\mathbb{R}^n))$ scaled by κ ,

$$U_i(x, \tau_i) = u(x, \theta_i) - \kappa \cdot c(\tau_i). \quad (9)$$

where

$$u(x, \theta) = -(x - \theta)^\top (x - \theta).$$

Voters ideal points are independent and the prior is spherical with mean normalized to the origin. We also assume that the prior is log-concave, which we use to show that voter learning benefits the moderate party (Proposition 3. The information cost is the expected reduction of variance, from the prior μ to the posterior π ,

$$c(\tau) = \mathbb{E}_\tau[\text{Var}_\mu(\theta) - \text{Var}_\pi(\theta)]$$

where $\text{Var}_\pi(X) := \mathbb{E}_\pi[\|\theta - \mathbb{E}_\pi[\theta]\|^2]$ is the multidimensional analogue of variance. This information cost is used in recent contributions (e.g., Ravid, Roesler, and Szentes, 2022; Thereze, 2022) and belongs to the class of posterior-mean separable information-costs axiomatized in Mensch and Malik (2024). The information cost is tractable because it is equivalent to the variance of the posterior means by the law of iterated variance,

$$\mathbb{E}_\tau[\text{Var}_\mu(\theta) - \text{Var}_\pi(\theta)] = \text{Var}_\tau(\mathbb{E}_\pi[\theta]).$$

Together with the quadratic policy utility, this information cost allows us to express the voter's objective, up to constants, as a simple function of the distribution over posterior means $\rho \in \Delta(\mathbb{R}^n)$, induced by $\tau \in \Delta(\Delta(\mathbb{R}^n))$:

$$\mathbb{E}_{\theta \sim \rho} \left[\max \left\{ \left\langle x_b - x_a, \theta - \frac{x_a + x_b}{2} \right\rangle, - \left\langle x_b - x_a, \theta - \frac{x_a + x_b}{2} \right\rangle \right\} - \kappa \langle \theta, \theta \rangle \right] \quad (10)$$

By Strassen's theorem, a distribution $\rho \in \Delta(\mathbb{R}^n)$ over posterior means is induced by some signal structure if and only if ρ is a mean-preserving contraction of the prior μ , $\rho \leq_{\text{MPS}} \mu$.³⁰ So, voters choose $\rho \in \Delta(\mathbb{R}^n)$ to maximize (10) subject to $\rho \leq_{\text{MPS}} \mu$. Due to the simple form of this objective, consisting of a piecewise linear and a quadratic function, we can obtain a closed-form solution for the voter's learning problem in some parameter region. While the qualitative results on voter learning hold for a more general class of information costs, this cost function proves particularly tractable to solve for endogenous party positions.

³⁰A distribution $\rho \in \Delta(\mathbb{R}^n)$ is a mean-preserving contraction of another distribution $\mu \in \Delta(\mathbb{R}^n)$ if there exists $\mathbb{R}^n$ -valued random variables X, Y such that $X \sim \rho$, $Y \sim \mu$, and $\mathbb{E}[Y|X] = X$.

Parties We assume that parties are policy-motivated (Wittman, 1973; Calvert, 1985). That is, the payoffs U_a and U_b of parties a and b , respectively, are

$$\begin{aligned} U_a(x_a, x_b) &= P_a(x_a, x_b)u(x_a, x_a^*) + (1 - P_a(x_a, x_b))u(x_b, x_a^*) \\ U_b(x_a, x_b) &= P_a(x_a, x_b)u(x_a, x_b^*) + (1 - P_a(x_a, x_b))u(x_b, x_b^*), \end{aligned}$$

where $P_a(x_a, x_b)$ is the probability that party a 's policy x_a gets implemented. We make the common assumption that this implementation probability is the vote share (Wittman, 1983; Wittman, 1990; Callander and Carbajal, 2022; Yuksel, 2022). A continuous mapping from the vote share to the implementation probability can result from additional noise voters or random turnout of partisans (e.g., Feddersen and Pesendorfer, 1996), or from non-majoritarian institutions in parliament.³¹ Computations suggest that our results generalize to S-shaped mappings from vote share to implementation probability instead of linear mappings, however, an S-shaped mapping may lead to equilibrium multiplicity.

Equilibrium We focus on subgame perfect pure-strategy equilibria, in which both parties receive positive expected vote shares.³²

5.2 Voter Learning

Since the posterior variance cost is reflection invariant, Blackwell monotonic, and posterior separable, our results on issue alignment and polarization from section 3 apply. Moreover, we can obtain additional results. Under the posterior variance information cost, when the prior is dispersed enough (see Appendix B for details), the two revealed ideal points $\theta_a, \theta_b \in \mathbb{R}^n$ acquired by voters have closed-form solutions, namely

$$\begin{aligned} \theta_a(x_a, x_b) &:= \frac{\langle \Delta x, \frac{x_a + x_b}{2} \rangle}{\langle \Delta x, \Delta x \rangle} \Delta x - \frac{\Delta x}{2\kappa} \\ \theta_b(x_a, x_b) &:= \frac{\langle \Delta x, \frac{x_a + x_b}{2} \rangle}{\langle \Delta x, \Delta x \rangle} \Delta x + \frac{\Delta x}{2\kappa} \end{aligned} \tag{11}$$

where $\Delta x = x_b - x_a$ is the party difference. Both revealed ideal points are on the line through the origin with direction equal to the party difference. The first term of (11) is the orthogonal projection of the party midpoint $\frac{x_a + x_b}{2}$ on said line. Both voter positions θ_a and θ_b are equally far from this projection. These closed-form solutions prove useful to characterize party positions, because they imply a simple expression for the vote share,

$$P_a(x_a, x_b) = \frac{1}{2} + \frac{\kappa}{2} \frac{\|x_b\|^2 - \|x_a\|^2}{\|x_b - x_a\|^2}.$$

³¹As Yuksel (2022) notes, if half the population consists of noise voters and their vote share for party a is uniformly distributed, then the implementation probability is exactly the vote share from non-noise voters.

³²There exists a disconnected class of equilibria, where voters acquire no information and all vote for one party. These equilibria can be easily ruled out, for example by introducing an arbitrarily small office benefit.

What complicates the analysis is that these closed-form solutions do not hold when there is no distribution ρ over posterior means θ_a and θ_b that is a mean-preserving contraction of the prior μ . Nevertheless, key properties of the solution can be obtained, such as the following proposition, which turns out to be crucial for understanding party positioning.

Costly voter learning generates a “bias” toward the moderate party: the moderate party receives more votes than the true share of voters whose ideal points are closer to its position. This bias emerges even though voters’ belief updating is unbiased in the Bayesian sense.

Proposition 3 (Bias toward Moderate Party). *Under $\|x_a\| < \|x_b\|$, party a obtains a higher expected vote share P_a than the true share of voters that are closer to x_a than to x_b . The expected vote share of party a decreases as the information cost parameter κ decreases or as party polarization $\|x_b - x_a\|$ increases, holding $\frac{x_a + x_b}{2}$ constant.*

The high-level intuition for the bias toward the moderate party is that asymmetric signals are less informative and hence cheaper to acquire. In other words, voters acquire biased signals to economize on their information cost. The bias toward the moderate party is smaller when information is cheaper or when parties are more polarized, because the motive to economize on the information cost is less important in these cases.

To unpack this intuition, it is helpful to consider how the optimal signal structure changes as the information cost parameter κ increases. When information is cheap, $\kappa \approx 0$, voters choose a threshold signal structure with a threshold approximating $(x_a + x_b)/2$, which perfectly separates voters according to whether they prefer x_a or x_b . Hence, the expected vote shares match the true shares of voters whose ideal point lie closer to each party. As κ increases, the optimal threshold becomes more extreme than $(x_a + x_b)/2$ (see Proposition 4), leading to a bias toward the moderate party. This happens because a more extreme threshold makes the binary signal more asymmetric, which—under a log-concave prior—reduces the informativeness and thus the cost of the signal. As information costs rise further, voters shift from adjusting the threshold to garbling the signal, using noise to reduce the information cost. Voters garble the signal structure in an asymmetric way, which further increases the expected vote share of the moderate party. Namely, the signal is more often distorted toward the moderate party. This is due a new consideration: if voters must distort the signal, distorting it toward the moderate party entails a smaller expected-utility loss than distorting it toward the extreme party. The reason is that a mistaken vote for the moderate party is typically closer to the voter’s ideal point than a mistaken vote for the extreme party.

Our notion of bias (toward the moderate party) is consistent with the definition of *media bias* given by Gentzkow, Shapiro, and Stone (2015). This connects our model to the literature on demand-driven media bias. Recall that one of the interpretations of voter learning is that information is provided by an intermediary such as news media. Whereas most demand-driven explanations of media bias involve psychological utility of news, we complement Che and Mierendorff (2019) by showing that media bias can be rational for voters to economize on their information cost. In their dynamic model, voters choose between left- and right-biased news subject to a time-based cost. For extreme beliefs, voters acquire news biased toward their current belief. By contrast, our model

allows for *flexible* information acquisition, so voters may acquire unbiased news if they want to. However, the acquired information is always biased toward the party closer to their initial belief.

Corollary 4 has important implications for electoral competition. The corollary shows that endogenous information creates two new forces affecting party positions: a **moderation force** and a **differentiation force**. First, the bias toward the moderate party creates an additional motive for parties to *moderate*. This motive is stronger, the more costly is information. Second, the extreme party has an incentive to *differentiate*, because party differentiation decreases the bias toward the moderate party. The flip-side of this is that the moderate party has an incentive to move closer to the extreme party, because this increases the bias toward itself.

We show in the following that the moderation motive is the central force at play in symmetric equilibria and the differentiation motive is the central for understanding asymmetric equilibria.

5.3 Equilibrium

5.3.1 Symmetric Setup

We first characterize the equilibrium under a symmetric setup, that is, when parties are equally ideologically extreme from the viewpoint of voters, $\|x_a^*\| = \|x_b^*\|$. For the statement of the theorem, define $\underline{\kappa}$ via

$$\frac{1 - \kappa}{\kappa} \|x_a^*\| = \mathbb{E}[\|\theta\|].$$

The parameter $\underline{\kappa}$ is the smallest information cost parameter for which the closed-form solution (11) for the voter learning problem still holds.

Theorem 5. *Let $\|x_a^*\| = \|x_b^*\|$ and $x_a^* \neq x_b^*$. The following constitutes an equilibrium.*

	Party platforms	Revealed voter ideology
$\kappa \in (\underline{\kappa}, 1)$:	Polarizing equilibrium $(x_a, x_b) = (1 - \kappa)(x_a^*, x_b^*)$	$\rho = \frac{1}{2}\delta\left(\frac{x_a}{\kappa}\right) + \frac{1}{2}\delta\left(\frac{x_b}{\kappa}\right)$
$\kappa \geq 1$:	Downsian equilibrium $(x_a, x_b) = (0, 0)$	$\rho = \delta(0)$

The equilibrium is unique if $\|x_b^ - x_a^*\| \leq \kappa \mathbb{E}[\|\theta\|]$ and $\kappa \neq 1/2$.*

The theorem describes two types of equilibria, depending on the size of the information cost.

For large information costs, there is a *Downsian equilibrium*, reminiscent of the median-voter theorem (Downs, 1957). In this equilibrium, parties fully converge at the voter's prior expectation of the optimal policy. Voters acquire no information and split their votes arbitrarily between the two parties.

For smaller information costs, there is what we call a *polarizing equilibrium*. In this equilibrium, parties diverge and voters acquire information, which polarizes the population into two groups. Surprisingly, the model allows for a simple closed-form solution of the party and voter positions. The smaller the information cost, the more parties and voters polarize.

The model features the same comparative statics—cheaper information increases polarization—as the model in section 4, where parties choose positions after voters learn. The comparative statics under the present timing are driven by the novel **moderation motive**. For a high-level intuition, recall Corollary 4: to economize on the cost of information, voter learning is biased toward the moderate party. This bias encourages parties to moderate. As information becomes cheaper, the motive to reduce the information cost weakens, decreasing the bias and leading to more party polarization.

To explain the intuition in more detail, let us first recall parties' trade-off when choosing positions. Policy-motivated parties trade off votes (winning often) and ideology (winning big). Parties gain votes by moving closer to voters' prior expectation. Parties gain ideological utility by moving closer to their own ideal policies, creating polarization. The optimal party positions equate the loss of votes with the gain in ideological utility from polarizing further. As a consequence, the smaller the loss of votes, the larger is the equilibrium level of polarization. Now, the bias toward the moderate party creates an additional loss of votes from polarizing further than the other party, dampening polarization. However, this bias toward the moderate party recedes as information becomes cheaper, creating more party polarization.

5.3.2 Asymmetric Setup

We now turn to the equilibrium when one party is more ideologically extreme, which we assume without loss is party b , so $\|x_a^*\| \leq \|x_b^*\|$.

To state the theorem, we define proj_v as the scalar projection on $v \in \mathbb{R}^n$,

$$\begin{aligned} \text{proj}_v : \mathbb{R}^n &\rightarrow \mathbb{R} \\ x &\mapsto \frac{\langle v, x \rangle}{\|v\|}, \end{aligned} \tag{12}$$

and $P_{y,z}$ as the orthogonal projection on the line through $y, z \in \mathbb{R}^n$,

$$\begin{aligned} P_{y,z} : \mathbb{R}^n &\rightarrow \mathbb{R}^n \\ x &\mapsto y + \frac{\langle z - y, x \rangle}{\|z - y\|} \frac{z - y}{\|z - y\|}. \end{aligned} \tag{13}$$

We prove the following theorem for the case that parties are ideologically “on different sides” of the center of the voter distribution. Formally, the theorem assumes that the scalar projections of party ideal policies x_a^* and x_b^* on the party ideological difference $x_b^* - x_a^*$ are on different sides of 0, that is, $\text{proj}_{x_b^* - x_a^*}(x_a^*) < 0 < \text{proj}_{x_b^* - x_a^*}(x_b^*)$.

Theorem 6. *Suppose $\text{proj}_{x_b^* - x_a^*}(x_a^*) < 0 < \text{proj}_{x_b^* - x_a^*}(x_b^*)$ and $\|x_a^*\| \leq \|x_b^*\|$. The positions*

$$\begin{aligned} x_a(\kappa) &= P_{(1-\kappa)x_a^*, (1-\kappa)x_b^*} \left(\frac{1-\kappa}{1-2\kappa} \left((1-\kappa)x_a^* + \kappa x_b^* \right) \right) \\ x_b(\kappa) &= P_{(1-\kappa)x_a^*, (1-\kappa)x_b^*} \left(\frac{1-\kappa}{1-2\kappa} \left((1-\kappa)x_b^* + \kappa x_a^* \right) \right) \end{aligned} \tag{14}$$

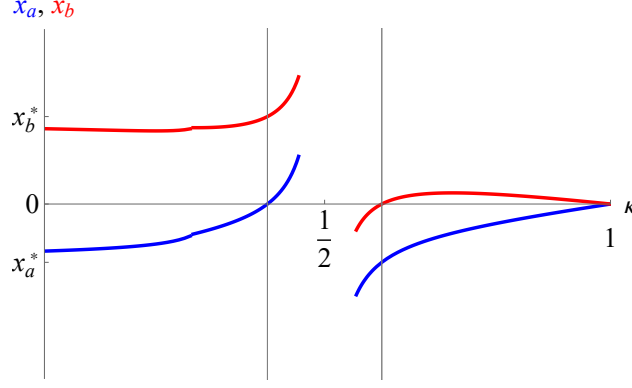


Figure 4: Party positions as a function of the information cost parameter κ

constitute an equilibrium for

$$\kappa \in \left(0, \frac{\hat{x}_b^* - \hat{x}_a^*}{3\hat{x}_b^* - \hat{x}_a^*}\right) \cup \left(\frac{\hat{x}_b^* - \hat{x}_a^*}{\hat{x}_b^* - 3\hat{x}_a^*}, 1\right) \quad \text{and} \quad \max\{\|x_a(\kappa)\|, \|x_b(\kappa)\|\} \leq \kappa \mathbb{E}[\|\theta\|]. \quad (15)$$

This is the unique equilibrium for κ satisfying (15) if additionally $\|x_b^* - x_a^*\| \leq \kappa \mathbb{E}[\|\theta\|]$.

If $\kappa \geq 1$, the unique equilibrium is $(x_a, x_b) = (0, 0)$.

Theorem 6 states that in equilibrium, party positions x_a and x_b are on the line through $(1 - \kappa)x_a^*$ and $(1 - \kappa)x_b^*$, for all κ satisfying (15). The positions of parties on said line are given by (14). As in the equilibrium under a symmetric setup, party polarization is given by the simple formula

$$x_b(\kappa) - x_a(\kappa) = (1 - \kappa)(x_b^* - x_a^*).$$

In particular, party polarization increases as the information cost κ decreases.

Figure 4 plots the equilibrium positions (x_a, x_b) as a function of κ under a one-dimensional policy space. For large information costs, $\kappa \geq 1$, parties converge fully as in the symmetric setup. For smaller information costs, $\kappa < 1$, parties polarize as κ decreases. For an intermediate range of κ , there is no equilibrium in pure strategies where both parties obtain positive expected vote shares. We comment on this below. For small κ , Figure 4 plots numerically solved equilibrium positions.

A notable feature of the equilibrium is that for intermediate information costs, both parties are on the same side of zero and one party is more extreme than its ideal policy. This contrasts with the Calvert-Wittman model under exogenous voter positions, where party positions necessarily lie between their ideal policies (Roemer, 1997).

This feature of the equilibrium is driven by the **differentiation motive**, which creates *chase-and-evade incentives*, similar to the literature on valence advantage (Grosseclose, 2001; Aragonés and Palfrey, 2002; Bernhardt, Buisseret, and Hidir, 2020). By the logic of the differentiation motive, in an asymmetric equilibrium, the extreme party has an incentive to differentiate (“evade”) and the moderate party has an incentive to move closer to (“chase”) the extreme party. When the information cost parameter κ is close to $1/2$, the differentiation motive become strong enough to

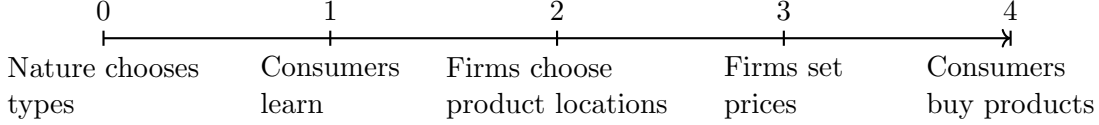
push the extreme party beyond its own ideal policy. In other words, endogenous voter learning gives parties an incentive to strategically extremize. The moderate party, on the other hand, moves so close to the extreme party that it is pulled beyond the center of the voter distribution, zero. If κ is very close to $1/2$, these chase-and-evade incentives become so strong that a pure-strategy equilibrium fails to exist, as is common in the literature on valence advantage (Aragones and Palfrey, 2002; Hummel, 2010).

6 Horizontally Differentiated Goods: Rising Markups

In this section, we adapt our model to a market context. We examine firms that produce horizontally differentiated products in some attribute space. A key difference between markets and politics is that consumption is private in markets, whereas policy is public: consumers can purchase the product that best fits their preferences, but while voters can support their preferred policy platform, ultimately only one policy is implemented for everyone. This distinction implies a potential benefit for product differentiation in markets—allowing consumers to select their optimal match—that is not present in politics. However, we show that even with this advantage, a lower cost of information still harms consumers. The reason is twofold: firms differentiate their products excessively from a social welfare standpoint, and unlike political parties, they also set prices. As consumers become more polarized due to lower information costs, firms can exploit their increased market power by raising prices, which further decreases consumer welfare. Methodologically, we show that endogenous consumer learning allows us to solve a multidimensional Hotelling model by reducing it to one dimension.

Our adaptation of the model can also be interpreted within a political-economy framework that includes valence competition. In this context, political parties not only choose policy platforms but also compete based on valence attributes such as competence of their candidates (see Ashworth and Bueno de Mesquita, 2009). Similar to how firms set prices in our industrial organization adaptation, parties may invest in valence to attract voters.

Adapting our model to a market context is largely a matter of reinterpretation. We reinterpret voters as consumers with unit demand, the policy space $\mathbb{R}^n$ as a product attribute space, and parties as firms. However, we need to make two important modifications to the model. First, we replace the exogenous valence shocks with prices, which are chosen by firms. Second, firms maximize profits rather than a combination of vote share and ideological utility, which our parties maximized. Apart from these two adjustments—the substitution of valence shocks with endogenous pricing and the change in the firms’ objectives—the models remain essentially the same, except for the previously mentioned welfare difference between private consumption and public policy. We also maintain the timing structure, adhering to the standard sequence in Hotelling models, where firms first choose product locations in attribute space and then set prices.



Consumers The utility U_i of consumer $i \in [0, 1]$ from consuming one unit of product x at price p given their preference θ and information cost $c(\tau)$ is

$$U_i(x, \theta, p, \tau) = u(x, \theta) - p - \kappa c(\tau).$$

We interpret $u(x, \theta)$ as the utility of consuming the good with attribute x . Consumers purchase whichever product gives them the higher expected utility.³³ We call the expected ideal point of a consumer their *revealed preference*.

Firms There are two firms, labelled a and b , in the market. Firms simultaneously choose their respective product locations, $x_a \in \mathbb{R}^n$ and $x_b \in \mathbb{R}^n$, and afterwards simultaneously choose their respective prices, p_a and p_b . Both firms have identical constant marginal costs, which we normalize to zero, so prices should be interpreted as markups. Firms maximize profits, that is market share times price. Given the distribution ρ of revealed preferences, the utility U_a of firm a is

$$U_a(x_a, x_b, p_a, p_b, \rho) = p_a \cdot \mathbb{E}_{\theta \sim \rho} [\mathbb{1}(u(x_a, \theta) - u(x_b, \theta) \geq p_a - p_b)],$$

and analogously for firm b .

We assume that consumers preferences are drawn from a normal distribution $\mathcal{N}(0, \Sigma)$ and consumers are restricted to normal signal structures. This is done, so the resulting distribution of revealed preferences is necessarily log-concave (it is normal), by which there exists a pure-strategy equilibrium of the price subgame for all product locations (Caplin and Nalebuff, 1991). We study the case of no aggregate uncertainty, assuming consumer ideal points are independently distributed, and maintain the focus on pure-strategy equilibria.

The following key lemma shows two important differences between the political and the market context. First, product locations are on the line of revealed preferences, whereas policy platforms are not necessarily on the line of revealed ideal points. Second, the direction of product differentiation is an eigenvector of ΣA , whereas the direction of platform differentiation is determined by the ideological difference between parties. Furthermore, the lemma allows us to reduce the model to one dimension and apply results from one-dimensional Hotelling models (Anderson, Goeree, and Ramer, 1997) for our subsequent equilibrium characterization.

Lemma 2. *Consumers' revealed preferences and product locations are supported on the same line, the direction of which is an eigenvector of ΣA .*

The intuition is as follows. Analogous to the political-economy context, consumers' best response to product locations is to learn such that revealed preferences are on a line with the direction

³³In other words, we assume that the utility from not purchasing any product is sufficiently low that all consumers prefer to buy one of the available products in equilibrium.

$\Sigma A(x_b - x_a)$. On the other hand, we show firms' best response is to locate their products on the line of consumer preferences: any other location is dominated by its projection on the line. Combining both, product differentiation $x_b - x_a$ must be parallel to $\Sigma A(x_b - x_a)$. That is, firms differentiate their products in a direction that is an eigenvector of ΣA .

Thus, there is a potential multiplicity of equilibria in this model. In fact, the following theorem shows that there is an equilibrium for any eigenvector of ΣA , provided the information cost parameter is small enough. To state the theorem, let $(v_1, \dots, v_n)$ denote a basis of A -normalized eigenvectors of ΣA , that is, there exists $\lambda_i \in \mathbb{R}$: $\Sigma A v_i = \lambda_i v_i$ and $v_i^\top A v_i = 1$ for all $i \in \{1, \dots, n\}$.³⁴ Let ρ denote the distribution of revealed consumer preferences in equilibrium.

Theorem 7. *The set of equilibria is fully characterized as follows:*

- *There is an equilibrium without learning, product differentiation, or markups:*

$$\sigma_\rho = 0, \quad x_a = x_b = 0, \quad p_a = p_b = 0.$$

- *For all $i \in \{1, \dots, n\}$, if*

$$\kappa < \frac{3}{2} v_i^\top A \Sigma A v_i,$$

there is an equilibrium with positive learning, product differentiation, and markups, given by

$$\rho = \mathcal{N}\left(0, \sigma_\rho^2 v_i v_i^\top\right), \quad \sigma_\rho^2 = v_i^\top A \Sigma A v_i - \frac{2}{3} \kappa, \quad -x_a = x_b = \frac{3}{4} \sqrt{2\pi} \sigma_\rho v_i, \quad p_a = p_b = 3\pi \sigma_\rho^2.$$

The intuition for the equilibrium without differentiation is as follows. In standard Hotelling models, firms differentiate their products to soften price competition (see, for example, d'Aspremont, Gabszewicz, and Thisse, 1979). However, this mechanism relies on consumer heterogeneity. If consumers do not acquire information and their revealed preferences are all located at zero, firms do not differentiate their products and charge no markups. Anticipating identical product locations, consumers choose not to learn, confirming this as an equilibrium.

In the equilibrium with consumer learning, consumers acquire information that disperses their revealed preferences ($\sigma_\rho > 0$). Firms respond by differentiating their products, which allows them to charge positive markups and earn profits. As the information cost κ decreases, consumers learn more, increasing σ_ρ . This greater consumer differentiation leads to higher product differentiation, further softening price competition and resulting in higher prices. Notably, prices increase quadratically in the dispersion of consumer preferences σ_ρ because *both* more differentiated products and more dispersed consumer preferences reduce competitive pressures.

A priori, it is unclear whether lower information costs benefit or harm consumers. While prices increase as κ decreases, product differentiation may benefit consumers by allowing a better match to products, and there is a direct positive effect from the reduced information cost. To address this question, we focus on the firm-optimal equilibrium, where products differentiate along the eigenvector v_i that maximizes $v_i^\top A \Sigma A v_i$. However, the comparative statics would be the same

³⁴Such a basis is given by the basis in which A and Σ^{-1} are simultaneously diagonalized as quadratic forms.

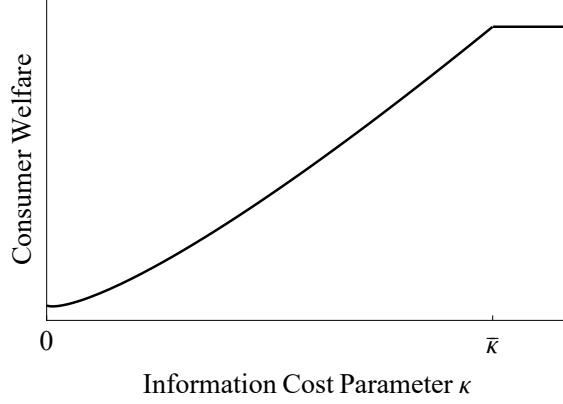


Figure 5: Consumer welfare increases as information becomes more costly.

if we considered the equilibrium associated with another eigenvector. We assess the aggregate utilitarian welfare of consumers, or, equivalently, their ex-ante utility. Define

$$\beta := 3e^{-(2+\frac{3}{4})\pi} \approx 0.04,$$

$$\bar{\kappa} := \frac{3}{2}v_i^\top A \Sigma A v_i.$$

Corollary 2 (Welfare Comparison). *As the information cost parameter κ increases, consumer welfare*

- *decreases strictly for $\kappa \in [0, \beta\bar{\kappa})$,*
- *increases strictly for $\kappa \in (\beta\bar{\kappa}, \bar{\kappa})$,*
- *remains constant for $\kappa \geq \bar{\kappa}$.*

Consumer welfare is maximized in the last case, in which there is no product differentiation.

Firm profits are decreasing in κ .

Figure 5 illustrates how consumer welfare varies with κ . Welfare decreases briefly when κ is small (from 0 to approximately $0.04 \bar{\kappa}$) and then increases until κ reaches $\bar{\kappa}$.

Our result shows that, except for a negligibly small region, cheaper information harms consumers. Although product differentiation can theoretically benefit consumers, firms differentiate excessively from a social standpoint (Anderson, Goeree, and Ramer, 1997), and increasingly so for lower information costs. Moreover, prices increase *quadratically* with σ_ρ because both more differentiated products and more dispersed consumer preferences reduce the intensity of competition. These negative effects outweigh the direct benefits consumers receive from lowering information costs. As a result, overall consumer welfare decreases when information becomes cheaper.

Our findings contribute to the literature on the welfare effects of information in markets. It is known that information can harm consumers by leading to higher prices (Moscarini and Ottaviani, 2001; Choi, Dai, and Kim, 2018; Armstrong and Zhou, 2022; Albrecht and Whitmeyer, 2023;

Biglaiser, Gu, and Li, 2024). Our analysis additionally incorporates endogenous product characteristics, which exacerbates the effect of consumer information on prices. New to this paper is, to the best of our knowledge, the comparative static result taking into account welfare effects of the cost of information, when both product characteristics and prices are endogenous.

7 Conclusion

In this paper, we dispense with the assumption that voters perfectly know their political preferences. Instead, voters can flexibly learn about their ideal points at a cost and do so for the purpose of expressing their political opinion in elections. Voters’ choice set shapes—through learning about ideal points—the revealed ideology in the population. Because voters’ choice set is constrained to the two party platforms in policy space, voters are not motivated to learn in a multidimensional or continuous way about their ideal points. As a result, revealed ideology displays issue alignment, and polarizes in the sense of approaching a binary distribution as valence becomes less uncertain. Voter learning predicts that polarization of voters and parties are mutually reinforcing and increase as information becomes cheaper. Finally, because voters only learn about the axis of party disagreement, policy is not responsive to dimensions of voter preferences that are orthogonal to this axis.

This paper opens several avenues for future research. One important question is what happens when voters learn jointly about their ideal points and party platforms. Such joint learning could explain the correlation between voters’ perceptions of candidate positions and their own ideologies (Hare, Armstrong, Bakker, Carroll, and Poole, 2015). Furthermore, while we considered voter learning occurring either before or after platform choice, in reality, both processes may coexist, especially since elections are held repeatedly: Some voters have acquired information during past elections, while other voters form opinions after observing current party campaigns. Additionally, examining repeated elections could shed more light on the dynamics of polarization.

References

- Abramowitz, A. (2010). *The disappearing center: Engaged citizens, polarization, and American democracy*. Yale University Press.
- Abramowitz, A., & Saunders, K. L. (2008). Is Polarization a Myth? *The Journal of Politics*, 70(2), 542–555.
- Acemoglu, D., Chernozhukov, V., & Woldz, M. (2016). Fragility of asymptotic agreement under Bayesian learning: Fragility of asymptotic agreement. *Theoretical Economics*, 11(1), 187–225.
- Achen, C., & Bartels, L. (2017). *Democracy for Realists: Why Elections Do Not Produce Responsive Government*. Princeton University Press.
- Albrecht, B. C., & Whitmeyer, M. (2023). Comparison Shopping: Learning Before Buying From Duopolists.
- Aliprantis, C. D., & Border, K. C. (2006). *Infinite Dimensional Analysis: A Hitchhikers Guide* (Vol. 10).

- Amari, S. (2016). *Information geometry and its applications* (Vol. 194). Springer.
- Anderson, S. P., Goeree, J. K., & Ramer, R. (1997). Location, location, location. *Journal of Economic Theory*, 77(1), 102–127.
- Ansolabehere, S., Rodden, J., & Snyder, J. M. (2008). The strength of issues: Using multiple measures to gauge preference stability, ideological constraint, and issue voting. *American Political Science Review*, 102(2), 215–232.
- Aragones, E., & Palfrey, T. R. (2002). Mixed Equilibrium in a Downsian Model with a Favored Candidate. *Journal of Economic Theory*, 103(1), 131–161.
- Armstrong, M., & Zhou, J. (2022). Consumer Information and the Limits to Competition. *American Economic Review*, 112(2), 534–577.
- Ashworth, S., & Bueno de Mesquita, E. (2009). Elections with platform and valence competition. *Games and Economic Behavior*, 67(1), 191–216.
- Bachmann, F., Sarasua, C., & Bernstein, A. (2024). Fast and Adaptive Questionnaires for Voting Advice Applications.
- Bafumi, J., & Herron, M. C. (2010). Leapfrog Representation and Extremism: A Study of American Voters and Their Members in Congress. *American Political Science Review*, 104(3), 519–542.
- Baliga, S., Hanany, E., & Klibanoff, P. (2013). Polarization and Ambiguity. *American Economic Review*, 103(7), 3071–3083.
- Bernhardt, D., Buisseret, P., & Hidir, S. (2020). The Race to the Base. *American Economic Review*, 110(3), 922–942.
- Biglaiser, G., Gu, J., & Li, F. (2024). Information Acquisition and Product Differentiation Perception. *American Economic Journal: Microeconomics*.
- Bischof, D., & Wagner, M. (2019). Do Voters Polarize When Radical Parties Enter Parliament? *American Journal of Political Science*, 63(4), 888–904.
- Bloedel, A. W., & Zhong, W. (2020). *The cost of optimally-acquired information* (tech. rep.). Technical report, Working paper, Stanford University.
- Bornschier, S. (2010a). Cleavage Politics and the Populist Right. The New Cultural Conflict in Western Europe. *The Social Logic of Politics*.
- Bornschier, S. (2010b). The New Cultural Divide and the Two-Dimensional Political Space in Western Europe. *West European Politics*, 33(3), 419–444.
- Callander, S., & Carbajal, J. C. (2022). Cause and Effect in Political Polarization: A Dynamic Analysis. *Journal of Political Economy*, 130(4), 825–880.
- Calvert, R. L. (1985). Robustness of the multidimensional voting model: Candidate motivations, uncertainty, and convergence. *American Journal of Political Science*, 69–95.
- Campbell, A., Converse, P. E., Miller, W. E., & Stokes, D. E. (1960). *The American Voter*. University of Chicago Press.
- Caplin, A., & Dean, M. (2013). *Behavioral implications of rational inattention with shannon entropy* (tech. rep.). National Bureau of Economic Research.
- Caplin, A., Dean, M., & Leahy, J. (2019). Rational inattention, optimal consideration sets, and stochastic choice. *The Review of Economic Studies*, 86(3), 1061–1094.
- Caplin, A., Dean, M., & Leahy, J. (2022). Rationally Inattentive Behavior: Characterizing and Generalizing Shannon Entropy. *Journal of Political Economy*, 130(6), 1676–1715.
- Caplin, A., & Nalebuff, B. (1991). Aggregation and imperfect competition: On the existence of equilibrium. *Econometrica: Journal of the Econometric Society*, 25–59.
- Carpini, M. X. D., & Keeter, S. (1996). *What Americans know about politics and why it matters*. Yale University Press.
- Casella, A. (2005). Storable votes. *Games and Economic Behavior*, 51(2), 391–419.

- Che, Y.-K., & Mierendorff, K. (2019). Optimal dynamic allocation of attention. *American Economic Review*, 109(8), 2993–3029.
- Choi, M., Dai, A. Y., & Kim, K. (2018). Consumer Search and Price Competition. *Econometrica*, 86(4), 1257–1281.
- Converse, P. (1964). The nature of belief systems in mass publics. *Ideology and Discontent*.
- Cunha, M., Osório, A., & Ribeiro, R. M. (2022). Endogenous Product Design and Quality When Consumers Have Heterogeneous Limited Attention.
- Curello, G., & Sinander, L. (2024). The comparative statics of persuasion.
- d’Aspremont, C., Gabszewicz, J. J., & Thisse, J.-F. (1979). On Hotelling’s” Stability in competition”. *Econometrica: Journal of the Econometric Society*, 1145–1150.
- DeMarzo, P. M., Vayanos, D., & Zwiebel, J. (2003). Persuasion bias, social influence, and unidimensional opinions. *The Quarterly journal of economics*, 118(3), 909–968.
- Downs, A. (1957). An Economic Theory of Political Action in a Democracy. *Journal of Political Economy*, 65(2), 135–150.
- Duggan, J., & Martinelli, C. (2011). A Spatial Theory of Media Slant and Voter Choice. *The Review of Economic Studies*, 78(2), 640–666.
- Duggan, J. (2017). A survey of equilibrium analysis in spatial models of elections. *Unpublished manuscript*.
- Dun, L., & Jessee, S. (2020). Demographic Moderation of Spatial Voting in Presidential Elections. *American Politics Research*, 48(6), 750–762.
- Duverger, M. (1954). *Political parties, their organization and activity in the modern state*. Methuen.
- Dworczak, P., & Kolotilin, A. (2023). The persuasion duality.
- Eguia, J., & Hu, T.-W. (2022). *Voter polarization and extremism* (tech. rep.).
- Enelow, J. M., & Hinich, M. J. (1989). A general probabilistic spatial theory of elections. *Public Choice*, 61(2), 101–113.
- Enke, B., Rodríguez-Padilla, R., & Zimmermann, F. (2023). Moral Universalism and the Structure of Ideology. *The Review of Economic Studies*, 90(4), 1934–1962.
- Esteban, J., & Ray, D. (1994). On the measurement of polarization. *Econometrica: Journal of the Econometric Society*, 819–851.
- Esteban, J., & Ray, D. (2012). Comparing Polarization Measures. In M. R. Garfinkel & S. Skaperdas (Eds.), *The Oxford Handbook of the Economics of Peace and Conflict*. Oxford University Press.
- Feddersen, T. J., & Pesendorfer, W. (1996). The swing voter’s curse. *The American economic review*, 408–424.
- Feddersen, T. J., & Sandroni, A. (2006). Ethical voters and costly information acquisition. *Quarterly Journal of Political Science*, 1(3), 287–312.
- Fiorina, M. P., & Abrams, S. J. (2008). Political polarization in the American public. *Annual Review of Political Science*, 11, 563.
- Fowler, A., Hill, S. J., Lewis, J. B., Tausanovitch, C., Vavreck, L., & Warshaw, C. (2022). Moderates. *American Political Science Review*, 1–18.
- Gentzkow, M., Shapiro, J. M., & Stone, D. F. (2015). Chapter 14 - Media Bias in the Marketplace: Theory. In S. P. Anderson, J. Waldfogel, & D. Strömberg (Eds.), *Handbook of Media Economics* (pp. 623–645). North-Holland.
- Groseclose, T. (2001). A Model of Candidate Location When One Candidate Has a Valence Advantage. *American Journal of Political Science*, 45(4), 862–886.
- Halff, H. M. (1976). Choice theories for differentially comparable alternatives. *Journal of Mathematical Psychology*, 14(3), 244–246.

- Hansson, I., & Stuart, C. (1984). Voting competitions with interested politicians: Platforms do not converge to the preferences of the median voter. *Public Choice*, 44(3), 431–441.
- Hare, C. (2022). Constrained Citizens? Ideological Structure and Conflict Extension in the US Electorate, 1980–2016. *British Journal of Political Science*, 52(4), 1602–1621.
- Hare, C., Armstrong, D. A., Bakker, R., Carroll, R., & Poole, K. T. (2015). Using Bayesian Aldrich-McKelvey Scaling to Study Citizens’ Ideological Preferences and Perceptions. *American Journal of Political Science*, 59(3), 759–774.
- Hare, C., Highton, B., & Jones, B. (2023). Assessing the Structure of Policy Preferences: A Hard Test of the Low Dimensionality Hypothesis. *The Journal of Politics*.
- He, J., & Natenzon, P. (2024). Moderate utility. *American Economic Review: Insights*, 6(2), 176–195.
- Hébert, B., & Woodford, M. (2021). Neighborhood-Based Information Costs. *American Economic Review*, 111(10), 3225–3255.
- Hébert, B., & Woodford, M. (2023). Rational inattention when decisions take time. *Journal of Economic Theory*, 208, 105612.
- Hill, S. J., & Tausanovitch, C. (2015). A Disconnect in Representation? Comparison of Trends in Congressional and Public Polarization. *The Journal of Politics*, 77(4), 1058–1075.
- Hinich, M. J. (1977). Equilibrium in spatial voting: The median voter result is an artifact. *Journal of Economic Theory*, 16(2), 208–219.
- Hu, L., Li, A., & Segal, I. (2023). The Politics of Personalized News Aggregation. *Journal of Political Economy Microeconomics*, 1(3), 463–505.
- Hummel, P. (2010). On the nature of equilibria in a Downsian model with candidate valence. *Games and Economic Behavior*, 70(2), 425–445.
- Inglehart, R., & Klingemann, H.-D. (1976). Party identification, ideological preference and the left-right dimension among Western mass publics. *Party identification and beyond*, 243–273.
- Jessee, S. A. (2009). Spatial voting in the 2004 presidential election. *American Political Science Review*, 103(1), 59–81.
- Jessee, S. A. (2012). *Ideology and spatial voting in American elections*. Cambridge University Press.
- Jewitt, I. (2004). Notes on the ‘Shape’ of Distributions.
- Jung, J., Kim, J. H., Matějka, F., & Sims, C. A. (2019). Discrete Actions in Information-Constrained Decision Problems. *The Review of Economic Studies*, 86(6), 2643–2667.
- Kamenica, E., & Gentzkow, M. (2011). Bayesian persuasion. *American Economic Review*, 101(6), 2590–2615.
- Kartik, N., Lee, S., Liu, T., & Rappoport, D. (2022). *Beyond Unbounded Beliefs: How Preferences and Information Interplay in Social Learning* (tech. rep.).
- Kőszegi, B., & Matějka, F. (2020). Choice simplification: A theory of mental budgeting and naive diversification. *The Quarterly Journal of Economics*, 135(2), 1153–1207.
- Kriesi, H., Grande, E., Lachat, R., Dolezal, M., Bornschier, S., & Frey, T. (2006). Globalization and the transformation of the national political space: Six European countries compared. *European Journal of Political Research*, 45(6), 921–956.
- Kriesi, H., Grande, E., Lachat, R., Dolezal, M., Bornschier, S., & Frey, T. (2008). *West European Politics in the Age of Globalization*. Cambridge University Press.
- Ladha, K. (1991). A spatial model of legislative voting with perceptual error. *Public Choice*, 68(1-3).
- Lauderdale, B. E. (2013). Does Inattention to Political Debate Explain the Polarization Gap between the U.S. Congress and Public? *Public Opinion Quarterly*, 77(S1), 2–23.
- Ledyard, J. O. (1984). The pure theory of large two-candidate elections. *Public Choice*, 44(1), 7–41.
- Leeper, T. J., & Slothuus, R. (2014). Political Parties, Motivated Reasoning, and Public Opinion Formation. *Political Psychology*, 35(S1), 129–156.

- Lelkes, Y. (2016). Mass polarization: Manifestations and measurements. *Public Opinion Quarterly*, 80(S1), 392–410.
- Lenz, G. S. (2012). *Follow the Leader?: How Voters Respond to Politicians' Policies and Performance*. University of Chicago Press.
- Levy, G., & Razin, R. (2015). Correlation neglect, voting behavior, and information aggregation. *American Economic Review*, 105(4), 1634–1645.
- Li, A., & Hu, L. (2023). Electoral accountability and selection with personalized information aggregation. *Games and Economic Behavior*, 140, 296–315.
- Lindbeck, A., & Weibull, J. W. (1987). Balanced-budget redistribution as the outcome of political competition. *Public choice*, 52(3), 273–297.
- Lipnowski, E., & Ravid, D. (2023). Predicting Choice from Information Costs.
- Maćkowiak, B., Matějka, F., & Wiederholt, M. (2023). Rational inattention: A review. *Journal of Economic Literature*, 61(1), 226–273.
- Malamud, S., & Schrimpf, A. (2022). Persuasion by Dimension Reduction.
- Malka, A., Lelkes, Y., & Soto, C. J. (2019). Are Cultural and Economic Conservatism Positively Correlated? A Large-Scale Cross-National Test. *British Journal of Political Science*, 49(3), 1045–1069.
- Martin, G. J., & Yurukoglu, A. (2017). Bias in cable news: Persuasion and polarization. *American Economic Review*, 107(9), 2565–2599.
- Martinelli, C. (2001). Elections with Privately Informed Parties and Voters. *Public Choice*, 108(1), 147–167.
- Martinelli, C. (2006). Would rational voters acquire costly information? *Journal of Economic Theory*, 129(1), 225–251.
- Matějka, F., & Tabellini, G. (2021). Electoral Competition with Rationally Inattentive Voters. *Journal of the European Economic Association*, 19(3), 1899–1935.
- McCarty, N. (2011). *Measuring Legislative Preferences*. Oxford University Press.
- McCarty, N. (2019). *Polarization: What everyone needs to know*®. Oxford University Press.
- McCarty, N., Poole, K. T., & Rosenthal, H. (2016). *Polarized America: The dance of ideology and unequal riches*. mit Press.
- McCarty, N., Rodden, J., Shor, B., Tausanovitch, C., & Warshaw, C. (2019). Geography, uncertainty, and polarization. *Political Science Research and Methods*, 7(4), 775–794.
- McMurray, J. C. (2023). Why the Political World is Flat: An Endogenous “Left” and “Right” in Multidimensional Elections.
- Mensch, J., & Malik, K. (2024). Posterior-Mean Separable Costs of Information Acquisition.
- Milgrom, P., & Shannon, C. (1994). Monotone Comparative Statics. *Econometrica*, 62(1), 157–180.
- Morris, S., & Strack, P. (2019). The wald problem and the relation of sequential sampling and ex-ante information costs. *Available at SSRN 2991567*.
- Moscarini, G., & Ottaviani, M. (2001). Price Competition for an Informed Buyer. *Journal of Economic Theory*, 101(2), 457–493.
- Mu, X., Pomatto, L., Strack, P., & Tamuz, O. (2021). From Blackwell Dominance in Large Samples to Rényi Divergences and Back Again. *Econometrica*, 89(1), 475–506.
- Nimark, K. P., & Sundaresan, S. (2019). Inattention and belief polarization. *Journal of Economic Theory*, 180, 203–228.
- Novák, V., Matveenko, A., & Ravaioli, S. (2024). The Status Quo and Belief Polarization of Inattentive Agents: Theory and Experiment. *American Economic Journal: Microeconomics*.
- Ortoleva, P., & Snowberg, E. (2015). Overconfidence in Political Behavior. *American Economic Review*, 105(2), 504–535.

- Palfrey, T. R., & Poole, K. T. (1987). The relationship between information, ideology, and voting behavior. *American journal of political science*, 511–530.
- Patty, J. W. (2002). Equivalence of Objectives in Two Candidate Elections. *Public Choice*, 112(1/2), 151–166.
- Patty, J. W. (2005). Local equilibrium equivalence in probabilistic voting models. *Games and Economic Behavior*, 51(2), 523–536.
- Perego, J., & Yuksel, S. (2022). Media Competition and Social Disagreement. *Econometrica*, 90(1), 223–265.
- Persson, T., & Tabellini, G. (2002). *Political economics: Explaining economic policy*. MIT press.
- Pew Reseach Center, W. D. (2014). Political polarization in the american public.
- Pomatto, L., Strack, P., & Tamuz, O. (2023). The Cost of Information: The Case of Constant Marginal Costs. *American Economic Review*, 113(5), 1360–1393.
- Poole, K. T., & Rosenthal, H. (1985). A Spatial Model for Legislative Roll Call Analysis. *American Journal of Political Science*, 29(2), 357–384.
- Posner, E. (1975). Random coding strategies for minimum entropy. *IEEE Transactions on Information Theory*, 21(4), 388–391.
- Ravid, D., Roesler, A.-K., & Szentes, B. (2022). Learning before trading: On the inefficiency of ignoring free information. *Journal of Political Economy*, 130(2), 346–387.
- Rayo, L., & Segal, I. (2010). Optimal information disclosure. *Journal of political Economy*, 118(5), 949–987.
- Roemer, J. E. (1997). Political–economic equilibrium when parties represent constituents: The unidimensional case. *Social Choice and Welfare*, 14(4), 479–502.
- Roemer, J. (1994). A theory of policy differentiation in single issue electoral politics. *Social Choice and Welfare*, 11(4).
- Shaked, M., & Shanthikumar, J. G. (Eds.). (2007). *Stochastic Orders*. Springer New York.
- Shor, B., & Rogowski, J. C. (2018). Ideology and the US Congressional Vote. *Political Science Research and Methods*, 6(2), 323–341.
- Sims, C. A. (2003). Implications of rational inattention. *Journal of Monetary Economics*, 50(3), 665–690.
- Smithies, A. (1941). Optimum Location in Spatial Competition. *Journal of Political Economy*, 49(3), 423–439.
- Spector, D. (2000). Rational debate and one-dimensional conflict. *The Quarterly Journal of Economics*, 115(1), 181–200.
- Sunstein, C. (2018). *#Republic: Divided Democracy in the Age of Social Media*. Princeton University Press.
- Tausanovitch, C., & Warshaw, C. (2012). *How Should We Choose Survey Questions to Measure Citizens’ Policy Preferences?* (Tech. rep.). Citeseer.
- Thereze, J. (2022). Adverse Selection and Endogenous Information.
- Tian, G., & Zhou, J. (1992). The maximum theorem and the existence of Nash equilibrium of (generalized) games without lower semicontinuties. *Journal of Mathematical Analysis and Applications*, 166(2), 351–364.
- Uhlig, H. (1996). A law of large numbers for large economies. *Economic Theory*, 8, 41–50.
- Van Der Brug, W., & Van Spanje, J. (2009). Immigration, Europe and the ‘new’ cultural dimension. *European Journal of Political Research*, 48(3), 309–334.
- Villani, C. (2009). *Optimal Transport* (M. Berger, B. Eckmann, P. De La Harpe, F. Hirzebruch, N. Hitchin, L. Hörmander, A. Kupiainen, G. Lebeau, M. Ratner, D. Serre, Y. G. Sinai, N. J. A. Sloane, A. M. Vershik, & M. Waldschmidt, Eds.; Vol. 338). Springer Berlin Heidelberg.

- Wittman, D. (1973). Parties as utility maximizers. *American Political Science Review*, 67(2), 490–498.
- Wittman, D. (1983). Candidate Motivation: A Synthesis of Alternative Theories. *The American Political Science Review*, 77(1), 142–157.
- Wittman, D. (1990). Spatial strategies when candidates have policy preferences. *Advances in the spatial theory of voting*, 66–98.
- Yoder, N. (2022). Designing incentives for heterogeneous researchers. *Journal of Political Economy*, 130(8), 2018–2054.
- Yuksel, S. (2022). Specialized Learning and Political Polarization. *International Economic Review*, 63(1), 457–474.
- Zaller, J. (1992). *The nature and origins of mass opinion*. Cambridge University Press.

A Appendix: Main Proofs

Throughout we use the notation $\langle x, y \rangle_A := x^\top A y$ and $\langle x, y \rangle := x^\top y$ for $x, y \in \mathbb{R}^n$, $A \in \mathbb{R}^{n \times n}$.

A.1 Theorem 1

Proof. The proof consists of two parts. In the first part, we show that only the A -projection³⁵ of the posterior mean on the platform difference $x_b - x_a$ is payoff relevant. In the second part of the proof, we show via a reflection argument that a voter acquires a distribution over posteriors such that the distribution over posteriors means has support on the line through the origin and $\Sigma A(x_b - x_a)$.

Part I The instrumental utility of τ , that is the objective of (P) neglecting the information cost, can be rewritten as follows:

$$\begin{aligned}
& \mathbb{E}_\tau \left[\mathbb{E}_\nu \left[\max \left\{ \mathbb{E}_\pi \left[-\langle \theta - x_a, \theta - x_a \rangle_A \right], \mathbb{E}_\pi \left[-\langle \theta - x_b, \theta - x_b \rangle_A \right] + \nu \right\} \right] \right] \\
&= \mathbb{E}_\tau \left[\mathbb{E}_\nu \left[\max \left\{ \mathbb{E}_\pi \left[-\langle x_a, x_a \rangle_A + 2\langle x_a, \theta \rangle_A - \langle \theta, \theta \rangle_A \right], \mathbb{E}_\pi \left[-\langle x_b, x_b \rangle_A + 2\langle x_b, \theta \rangle_A - \langle \theta, \theta \rangle_A \right] + \nu \right\} \right] \right] \\
&= \mathbb{E}_\tau \left[\mathbb{E}_\nu \left[\max \left\{ -\langle x_a, x_a \rangle_A + 2\langle x_a, \mathbb{E}_\pi[\theta] \rangle_A, -\langle x_b, x_b \rangle_A + 2\langle x_b, \mathbb{E}_\pi[\theta] \rangle_A + \nu \right\} \right] - \mathbb{E}_\pi \left[\langle \theta, \theta \rangle_A \right] \right] \\
&= \mathbb{E}_\tau \left[\mathbb{E}_\nu \left[\max \left\{ -\langle x_a, x_a \rangle_A + 2\langle x_a, \mathbb{E}_\pi[\theta] \rangle_A, -\langle x_b, x_b \rangle_A + 2\langle x_b, \mathbb{E}_\pi[\theta] \rangle_A + \nu \right\} \right] \right] - \mathbb{E}_\mu \left[\langle \theta, \theta \rangle_A \right] \\
&= \mathbb{E}_\tau \left[\mathbb{E}_\nu \left[\max \left\{ -\langle x_a, x_a \rangle_A + \langle x_a - x_b, \mathbb{E}_\pi[\theta] \rangle_A, -\langle x_b, x_b \rangle_A + \langle x_b - x_a, \mathbb{E}_\pi[\theta] \rangle_A + \nu \right\} \right] \right] + C_1 \\
&= \mathbb{E}_\tau \left[\mathbb{E}_\nu \left[\max \left\{ \left\langle x_a - x_b, \mathbb{E}_\pi[\theta] - \frac{x_a + x_b}{2} \right\rangle_A, \left\langle x_b - x_a, \mathbb{E}_\pi[\theta] - \frac{x_a + x_b}{2} \right\rangle_A + \nu \right\} \right] \right] + C_1 + C_2
\end{aligned} \tag{16}$$

where

$$\begin{aligned}
C_1 &= \mathbb{E}_\tau \left[\langle x_a + x_b, \mathbb{E}_\pi[\theta] \rangle_A \right] - \mathbb{E}_\mu \left[\langle \theta, \theta \rangle_A \right] = \mathbb{E}_\mu \left[\langle x_a + x_b, \theta \rangle_A - \langle \theta, \theta \rangle_A \right] \\
C_2 &= -\frac{1}{2} (\langle x_a, x_a \rangle_A + \langle x_b, x_b \rangle_A)
\end{aligned}$$

are constants. In the third line, we used that the expectation is a linear operator and that the inner product is linear. In the fourth line, we used the law of iterated expectations to show only the first moment of the posterior is payoff-relevant.

Because $\mathbb{E}_\nu[\max\{x, y + \nu\}]$ is a function of x , y , and the distribution of ν only, (16) shows that the instrumental value of information depends only on the distribution of the A -projection of the posterior mean $\mathbb{E}_\pi[\theta]$ on the platform difference $x_b - x_a$.

³⁵For any symmetric, positive definite matrix $B \in \mathbb{R}^{n \times n}$ and vector $v \in \mathbb{R}^n$, we refer to $\frac{\langle v, \theta \rangle_B}{\langle v, v \rangle_B} v = \frac{v^\top B \theta}{v^\top B v} v$ as the B -projection of θ on v .

Part II Suppose, for the sake of contradiction, that the voter acquires a distribution τ over posteriors that induces a distribution ρ of posterior means, which does not have support on the line through the prior mean (origin) and $\Sigma A(x_b - x_a)$. We construct—through reflection, mixing, and garbling—a strictly cheaper but instrumentally as valuable distribution $\hat{\tau}$ over posteriors that induces a distribution $\hat{\rho}$ over posterior means supported on the line through the origin and $\Sigma A(x_b - x_a)$. The three steps of our construction of $\hat{\tau}$ are visualized in Figure 2.

Step 1 - Reflection: We define a Bayes-consistent distribution $\text{Ref}(\tau)$ of the “reflected” posteriors that has the same information cost as well as instrumental value as τ .

To prove the result for general A and Σ , we make use of the Σ^{-1} -reflection Ref across the line through the origin and $\Delta\hat{x} := \Sigma A(x_b - x_a)$, defined as

$$\text{Ref}(\theta) = 2 \frac{\langle \Delta\hat{x}, \theta \rangle_{\Sigma^{-1}}}{\langle \Delta\hat{x}, \Delta\hat{x} \rangle_{\Sigma^{-1}}} \Delta\hat{x} - \theta.$$

This is a well-defined reflection because the symmetric, positive definite matrix Σ has a symmetric, positive definite inverse Σ^{-1} . In the simple case when Σ and A are equal to the identity matrix, Ref is simply the standard reflection across the line through the origin and $x_b - x_a$. In the general case, this reflection is useful for two reasons.

First, the Σ^{-1} -projection on $\Delta\hat{x}$ preserves the instrumental value. That is because the projection is equivalent to the A -projection on the platform difference $x_b - x_a$ by

$$\langle x_b - x_a, \theta \rangle_A = (x_b - x_a)^\top (\Sigma^{-1} \Sigma) A (x_b - x_a) = (\Sigma A (x_b - x_a))^\top \Sigma^{-1} \theta = \langle \Delta\hat{x}, \theta \rangle_{\Sigma^{-1}}. \quad (17)$$

That the reflection Ref preserves the Σ^{-1} -projection on $\Delta\hat{x}$ thus implies preserving the payoff-relevant A -projection on the platform difference $x_b - x_a$.

Second, Ref preserves the prior μ . Note that Ref is a linear function, which we can describe as multiplication by a matrix Q , $\text{Ref}(\theta) = Q\theta$. By Ref being a reflection, we have $Q = Q^{-1}$. Ref is a reflection with respect to inner product Σ^{-1} , so Ref preserves the distance induced by Σ^{-1} . Hence, we have $Q^\top \Sigma^{-1} Q = \Sigma^{-1}$, inverting which delivers $Q \Sigma Q^\top = \Sigma$. The characteristic function of $Q\theta$, as a random vector, satisfies for all $t \in \mathbb{R}^n$, $\Phi_{Q\theta}(t) = \Phi_\theta(Q^\top t) = \psi(t^\top Q \Sigma Q^\top t) = \psi(t^\top \Sigma t) = \Phi_\theta(t)$. Thus, $\text{Ref}(\theta) = Q\theta$ and θ have the same distribution.

The reflection Ref of the state space induces an according reflection on posteriors and on distributions over posteriors through the pushforward and iterated pushforward, which we both denote by Ref as well.³⁶ Intuitively, we are simply relabeling the states.

The distribution $\text{Ref}(\tau)$ is Bayes-consistent, since for all Borel sets $A \in \mathbb{R}^n$,

$$\int \pi(A) d(\text{Ref}(\tau)) = \int \text{Ref}(\pi)(A) d\tau = \text{Ref}(\mu)(A) = \mu(A),$$

where the last equality holds by the prior μ being invariant under Ref .

³⁶The pushforward $\text{Ref}_*: \Delta(\mathbb{R}^n) \rightarrow \Delta(\mathbb{R}^n)$ is formally defined via $\text{Ref}_*(\pi)(A) = \pi(\text{Ref}^{-1}(A))$ for Borel sets $A \subseteq \mathbb{R}^n$. The iterated pushforward on $\Delta(\Delta(\mathbb{R}^n))$ is simply $(\text{Ref}_*)^*$. For ease of reading, we write Ref for both Ref_* and $(\text{Ref}_*)^*$.

As argued above, Ref preserves the A -projection on $x_b - x_a$, and by linearity of Ref it maintains the distribution of the A -projection of the posterior *mean* on the platform difference. Thus, the instrumental value of τ is preserved under Ref . The information cost is also preserved under Ref since

$$\int D(\pi||\mu)d(\text{Ref}(\tau)) = \int D(\text{Ref}(\pi)||\mu)d\pi = \int D(\text{Ref}(\pi)||\text{Ref}(\mu))d\tau = \int D(\pi||\mu)d\tau,$$

where the last equality holds because the Kullback-Leibler divergence is invariant under coordinate transformations. Thus, the voter is indifferent between τ and $\text{Ref}(\tau)$. Finally, the distribution over posteriors means induced by $\text{Ref}(\tau)$ is simply $\text{Ref}(\rho)$, the reflection of the distribution over posterior means induced by τ .

Step 2- Mixing: It follows immediately that the mixture $\frac{1}{2}\tau + \frac{1}{2}\text{Ref}(\tau)$ is also Bayes-consistent and has the same instrumental value. It also has the same information cost by posterior separability of the information cost, which implies that the cost is linear under mixing. By posterior separability, the information cost can be written as an expectation with respect to the distribution τ over posteriors, which is linear in the distribution τ .

Step 3 - Garbling: Finally, we take a certain mean-preserving contraction of $\frac{1}{2}\tau + \frac{1}{2}\text{Ref}(\tau)$ to reach $\hat{\tau}$ (which corresponds to a garbling of the corresponding signal structure), which is also Bayes-consistent and has the same instrumental value, *but has a lower information cost*. We use the mean preserving contraction that contracts all posteriors whose means have the same Σ^{-1} -projection on the line through $\Sigma A(x_b - x_a)$. Any mean-preserving contraction is Bayes-consistent. The contraction $\hat{\tau}$ preserves the instrumental value of information since it preserves the distribution of the A -projection of the posterior mean on $x_b - x_a$. Crucially, the distribution over posterior means $\hat{\rho}$ induced by $\hat{\tau}$ has support on the line through the origin and $\Sigma A(x_b - x_a)$. The reason is that $\frac{1}{2}\tau + \frac{1}{2}\text{Ref}(\tau)$ is constructed to be symmetric around this line with respect to the Σ^{-1} -projection. Finally, a mean-preserving contraction lowers the information cost by convexity of the Kullback-Leibler divergence $D_{KL}(\pi||\mu)$ in its first argument. In fact, the Kullback-Leibler divergence is strictly convex for π that are absolutely continuous with respect to μ by strict convexity of $x \log x$ and $D_{KL}(\pi||\mu) = \int \log\left(\frac{d\pi}{d\mu}\right) \frac{d\pi}{d\mu} d\mu$. Without loss, the posteriors π induced by τ are almost surely absolutely continuous with respect to prior μ , otherwise τ has infinite cost and is clearly suboptimal. Thus, $\hat{\tau}$ has a strictly lower information cost than the original distribution τ if τ did not have posterior means already on the line, in which case the mean-preserving contraction is strict. $\square$

In Appendix D.3, we discuss more general information costs, such as certain distance-based costs, for which this proof works.

A.2 Theorem 2

Proof. We show, for completeness, that *any* optimal signal structure has no more signals than actions under degenerate valence shock, $\nu \sim \delta_0$. The proof uses only convexity properties of the Kullback-Leibler divergence D_{KL} .

The voter acquires a distribution τ over posteriors π , such that the posterior is almost surely absolutely continuous with respect to the prior μ . Otherwise, the voter obtains a negative infinite payoff and could do better by acquiring no information. Suppose, for the sake of contradiction, that the voter acquires a signal structure such that the induced distribution τ over posteriors has a support with more than two posteriors. Then, we can strictly improve the voter's utility by garbling the signal structure based on the action recommendation (after resolving indifferences between parties a and b for a , say). More precisely, we partition the space of posteriors $\Delta(\mathbb{R}^n)$ into the subset Δ_a on which voting for a is weakly preferred and another subset Δ_b on which voting for b is strictly preferred. The garbling corresponds to a mean-preserving contraction of τ , namely contracting all posteriors in Δ_a and in Δ_b , respectively, inducing a binary distribution over posteriors. The voter's utility from a distribution τ over posteriors π is the expected value of the value function, which consists of the instrumental value and the Kullback-Leibler divergence,

$$\int (\max \{ \mathbb{E}_\pi[u(x_a, \theta)], \mathbb{E}_\pi[u(x_b, \theta)] \} - \kappa D_{\text{KL}}(\pi || \mu)) d\tau.$$

On each Δ_a and Δ_b , the instrumental value (the max term) is linear in the posterior. The divergence D_{KL} is strictly convex in posterior π for π absolutely continuous with respect to μ , as argued in step 3 of the proof of Theorem 1. Thus, the value function is strictly concave on the support of τ in Δ_a and in Δ_b , each. By Jensen's inequality, our mean-preserving contraction on Δ_a and Δ_b weakly improves the voter's utility. Because the mean-preserving contraction is strict on at least one of Δ_a and Δ_b , the voter's utility improves strictly.

The second part of Theorem 2 follows immediately from Proposition 7 and Lemma 24 in the Appendix D. $\square$

A.3 Lemma 1

Proof. First, we show the best response x_a of party a to x_b and ρ necessarily satisfies the first-order condition of party a 's objective. The argument is analogous for party b . If party a chooses $x_a = x_a^*$, it obtains utility greater or equal to zero because x_a delivers zero ideological utility, $u(x_a, x_a^*) = 0$, and a non-negative vote share. If x_a is outside the ellipse described by $u(x_a, x_a^*) \geq -m$, then the resulting utility is negative, because the utility from the vote share can be at most m . Thus, all x_a outside this compact ellipse are suboptimal and by differentiability of the vote share in x_a , shown below, the maximum is obtained on this ellipse and necessarily satisfies the first-order condition, which we analyze next.

Taking the gradient ∇ with respect to x_a of the objective of party a , we obtain the necessary

first-order condition of the optimal platform x_a given ρ and x_b :

$$\begin{aligned}
& \nabla \left(m \int F_\nu(u(x_a, \theta) - u(x_b, \theta)) d\rho(\theta) + u(x_a, x_a^*) \right) = 0 \\
& \Leftrightarrow m \int \underbrace{f_\nu(u(x_a, \theta) - u(x_b, \theta))}_{=:w(\theta)} \nabla u(x_a, \theta) d\rho(\theta) + \nabla u(x_a, x_a^*) = 0 \\
& \Leftrightarrow \int mw(\theta) 2A(x_a - \theta) d\rho(\theta) + 2A(x_a - x_a^*) = 0 \\
& \Leftrightarrow 2A \left(m \int w(\theta)(x_a - \theta) d\rho(\theta) + (x_a - x_a^*) \right) = 0 \\
& \Leftrightarrow x_a = \frac{m \int w(\theta)\theta d\rho(\theta) + x_a^*}{m \int w(\theta) d\rho(\theta) + 1}
\end{aligned}$$

The integral of the expected vote share is well-defined. We can exchange integration and differentiation because the partial derivative of the integrand exists and is bounded in absolute value by an integrable function in θ . The latter holds because ν has finite first absolute moment and $\nabla u(x_a, \theta)$ is linear in θ . The last equivalence uses that A is symmetric and positive definite, so its kernel is $\{0\}$.

The result for platform x_b is analogous. Together, this implies

$$x_b - x_a = \frac{x_b^* - x_a^*}{m \int_{\mathbb{R}^n} w(\theta) d\rho(\theta) + 1}, \quad (18)$$

so $x_b - x_a$ is parallel to $x_b^* - x_a^*$. □

A.4 Theorem 3

Proof. By Lemma 25 in Appendix D, under the restriction to normal signal structures, it still holds that voters' revealed ideal points are on the line through the origin and $\Sigma A(x_b - x_a)$. The first-order conditions that characterize the equilibrium platforms (Lemma 1) are unaffected by the component of voter ideal points orthogonal (with respect to A) to $x_b - x_a$. That is because the ideal point θ enters the first-order conditions only via the utility difference $u(x_a, \theta) - u(x_b, \theta) = 2\lambda \langle x_a - x_b, \theta - \frac{x_a + x_b}{2} \rangle_A$. This utility difference is unaffected by the component of θ orthogonal (with respect to A) to $x_b - x_a$. Thus, while in the following proof we assume that the line of voter ideal points is parallel to $x_b - x_a$, all steps generalize to a line of voter ideal points that is slanted with respect to $x_b - x_a$. Furthermore, in any equilibrium, $x_b - x_a$ is parallel to $x_b^* - x_a^*$ by (18). To simplify exposition we change into an orthonormal basis of A in which $x_b^* - x_a^*$, and hence $x_b - x_a$, is parallel to the first basis vector. Such a basis exists by the Gram-Schmidt algorithm.

We show all equilibria are symmetric. By Lemma 1, equilibrium platforms can be written as a weighted average of voter ideal points and an *aggregate voter ideal point* $\bar{\theta}$,

$$\bar{\theta} := \frac{\int w(\theta)\theta d\rho(\theta)}{\int w(\theta) d\rho(\theta)},$$

by

$$x_j = \frac{m \int w(\theta) \theta d\rho(\theta) + x_j^*}{m \int w(\theta) d\rho(\theta) + 1} = \frac{m\bar{\theta} + \frac{1}{\int w(\theta) d\rho(\theta)} x_j^*}{m + \frac{1}{\int w(\theta) d\rho(\theta)}} \quad (19)$$

for $j = a, b$. We show in any equilibrium, the aggregate voter ideal point $\bar{\theta}$ is zero, which implies a symmetric equilibrium

$$(x_a, x_b) = \frac{\frac{1}{\int w(\theta) d\rho(\theta)}}{m + \frac{1}{\int w(\theta) d\rho(\theta)}} (x_a^*, x_b^*).$$

Suppose, for the sake of contradiction, that the aggregate voter ideal point $\bar{\theta}$ was not zero. By the paragraph above, the revealed voter ideal points are on a line parallel to the first basis vector. Thus, the aggregate voter ideal point $\bar{\theta}$ is on this line. Suppose without loss that its first component is positive, $\bar{\theta}_1 > 0$. We show this implies that the platform midpoint $\bar{x} := \frac{x_a + x_b}{2}$ must be to the left of the aggregate voter ideal point $\bar{\theta}$, which in turn must be to the left of the platform midpoint, creating a contradiction. By (19) and $x_{a,1}^* = -x_{b,1}^*$, the first component of the platform midpoint $\bar{x}_1$ is positive by

$$\bar{x}_1 = \frac{m\bar{\theta}_1 + \frac{1}{\int w(\theta) d\rho(\theta)} (x_{a,1}^* + x_{b,1}^*)}{m + \frac{1}{\int w(\theta) d\rho(\theta)}} = \frac{m}{m + \frac{1}{\int w(\theta) d\rho(\theta)}} \bar{\theta}_1 \in (0, \bar{\theta}_1). \quad (20)$$

Under normal signals, the distribution over posterior means is symmetric around 0 and quasi-concave. Thus, a positive party midpoint, $\bar{x}_1 > 0$, implies that the weighted mass of ideal points to the left of $\bar{x}$ is greater than to the right, so the aggregate voter ideal point must be to the left of the party midpoint. Formally, writing vectors in row-notation, $\forall y > 0$, $w((\bar{x}_1 - y, 0, \dots, 0)) = w((\bar{x}_1 + y, 0, \dots, 0))$ but the density of revealed voter ideal points is greater at $(\bar{x}_1 - y, 0, \dots, 0)$ than at $(\bar{x}_1 + y, 0, \dots, 0)$. Thus, $\bar{\theta}_1 = \int w(\theta) \theta d\rho(\theta) < \bar{x}_1$ in contradiction to (20).

Symmetric party platforms

$$(x_a, x_b) = \alpha(x_a^*, x_b^*) \quad (21)$$

satisfy the first-order conditions of optimality if

$$\alpha = \frac{1}{m \int f_\nu(\alpha \langle x_b^* - x_a^*, \theta \rangle) d\rho(\theta) + 1}. \quad (22)$$

We call α the degree of platform polarization. We establish equilibrium existence by showing that there exists a degree of platform polarization $\alpha \in (0, 1)$ that satisfies (22) where ρ is the induced distribution over revealed ideal points when the party platforms satisfy (21). We do so by constructing an equilibrium correspondence whose fixed point exists by monotonicity properties. Conceptually, this proof of pure-strategy equilibrium existence and the subsequent comparative statics result are similar to those in supermodular games.

We construct a correspondence G from $[0, 1]$ to $(0, 1)$ as a concatenation of two correspondences, g_1 and g_2 . Let g_1 map $\alpha \in [0, 1]$ to set of distributions ρ over posterior means induced by some optimal learning strategy given the degree of platform polarization α . We can restrict attention

to one-dimensional normal signals, which can be parametrized and ordered by the variance σ_ρ^2 of ρ . Because this variance is bounded by the prior variance in that dimension, standard arguments deliver that g_1 is nonempty compact-valued. Let g_2 map some distribution ρ of revealed ideal points to the set of equilibrium α that satisfy the first-order condition (22). The function g_2 is nonempty-valued with values in $(0, 1)$ because both the left-hand side and right-hand side of (22) are continuous and at $\alpha = 0$ the right-hand side is larger while at $\alpha = 1$, the left-hand side is larger. By the intermediate value theorem, a solution α exists. Define the correspondence $G = g_2 \circ g_1$,

$$\begin{aligned} G: [0, 1] &\longrightarrow 2^{[0, 1]} \\ \alpha &\xrightarrow{g_1} \{\rho\} \xrightarrow{g_2} \{\alpha\} \end{aligned}$$

Lemma 3. *Platform polarization increases voter polarization, that is $\min g_1$ and $\max g_1$ are strictly increasing.*

This lemma follows from Proposition 1.

Lemma 4. *Voter polarization increases platform polarization, that is, $\min g_2$ and $\max g_2$ are strictly increasing in the variance σ_ρ^2 of the symmetric, normal distribution ρ of voter ideal points.*

Proof. The smallest and largest α that solve (22) exist because g_2 is nonempty by the above and because the left- and right-hand side of (22) are continuous, so the preimage of $\{0\}$ under the continuous difference between the left- and right-hand side is closed.

A higher variance σ_ρ^2 of ρ implies that voters are strictly further away from 0 (the projection of the party midpoint) in first-order stochastic dominance. This implies that the term $\int f_\nu(\alpha \langle x_b^* - x_a^*, \theta \rangle) d\rho(\theta)$ strictly decreases by strict quasi-concavity of f_ν . Thus, the right-hand side of (22) strictly increases pointwise, which implies a greater smallest and largest α that solves (22). $\square$

Together, this implies that the minimum and maximum of $G : [0, 1] \rightarrow [0, 1]$,

$$\begin{aligned} (\min G)(\alpha) &:= \min\{G(\alpha)\} = \min g_2(\min g_1(\alpha)) \\ (\max G)(\alpha) &:= \max\{G(\alpha)\} = \max g_2(\max g_1(\alpha)) \end{aligned}$$

are strictly increasing. Fixed points of $\min G$ and $\max G$ correspond to equilibrium degrees of platform polarization. These exist by Tarski's fixed point theorem, viewing $[0, 1]$ with the usual ordering as a complete lattice using monotonicity of $\min G$ and $\max G$.

Finally, we show the comparative statics result.

Lemma 5. *A smaller κ implies that the smallest and largest fixed point of G weakly increase.*

Proof. A smallest and largest fixed point of G exist by Tarski's fixed point theorem.

The functions $\min G$ and $\max G$ weakly increase pointwise as κ decreases. This implies that the smallest and largest intersection of G with the identity function on $[0, 1]$ increase strictly, leading to a higher smallest and largest equilibrium α . To prove this, note first that g_2 is unchanged. The

functions $\min g_1$ and $\max g_1$ weakly increase pointwise by Proposition 1. Together with $\min g_2$ and $\max g_2$ being non-decreasing, this implies $\min G$ and $\max G$ are weakly higher pointwise. $\square$

This concludes the proof of Theorem 3. $\square$

A.5 Theorem 4

Before we prove Theorem 4, we formally describe the strategies and payoffs of players. We also pave the way for our proof by introducing a way to represent the extensive-form strategies of our two parties and continuum of voters as a static game between only four players.

Players and Strategies: Parties $j \in \{a, b\}$ choose their platforms conditional on the realized public opinion signal $s \in S$. Formally, a strategy of party j is a function $x_j : S \rightarrow \mathbb{R}^n$.

After learning, voters choose who to vote for conditional on the public signal s and the realized party platforms $(x_a(s), x_b(s))$. Formally, voters choose vote choice functions $v : S \times \mathbb{R}^n \times \mathbb{R}^n \rightarrow \{a, b\}$ that map $(s, x_a(s), x_b(s))$ into a choice among parties. Because the posterior mean is a sufficient statistic for optimal voting behavior (Remark 1), we can code subgame-perfect vote choice functions in the following strategically equivalent reduced-form way: voters choose a posterior mean $p \in \mathbb{R}^n$ conditional on each public signal s , incorporating that, subsequently, they choose optimally between $x_a(s)$ and $x_b(s)$ given posterior mean p and their valence shock. We denote a generic reduced-form strategy by $p_S \in (\mathbb{R}^n)^S$. As usual in rational inattention, it is without loss to identify the signal space with the set of actions, $\mathcal{S} = (\mathbb{R}^n)^S$. Thus, voter i 's extensive-form strategy is reduced to a signal structure (stochastic kernel) $\sigma_i : \Omega \times D \rightarrow \Delta(\mathbb{R}^n)^S$.

Because voters are ex-ante homogeneous and we assume that all voters acquire the same signal structure, we model our continuum of voters through a *representative voter*, who chooses a signal structure $\sigma : \Omega \times D \rightarrow (\mathbb{R}^n)^S$.

As mentioned in the main text, a single, infinitesimal voter cannot affect the realized signal distribution. Modeling voters as a representative voter, we have to ensure that, when we define payoffs, the representative voter's signal structure does not affect the public signal. Therefore, we introduce a fourth, *fictional player* who also chooses a signal structure $\sigma_f : \Omega \times D \rightarrow \Delta((\mathbb{R}^n)^S)$. We will define the payoff of the representative voter such that in equilibrium, it mimics the signal structure of the representative voter, ensuring consistency of the public opinion signal with voters' learning strategies.

In particular, the public signal is obtained as follows from σ_f . By the continuum of voters and the compact support of p_S , which implies finite second moments, we can apply a law of large numbers for a continuum of random variables (Uhlig, 1996) if we interpret realized population distributions as Pettis integrals. In particular, if voters acquire the signal structure σ_f , then, conditional on the aggregate state ω , the realized distribution of signals is deterministically $\sum_{\delta} \mu(\omega, \delta) \sigma_f(\cdot | \omega, \delta)$. Thus, the probability of a certain public opinion signal s conditional on aggregate state ω and

learning strategy σ_f is

$$\sigma_p(s|\omega, \sigma_f) := \sigma_p \left(s \middle| \sum_{\delta} \mu(\omega, \delta) \sigma_f(\cdot | \omega, \delta) \right).$$

From this definition, it follows that the probability $\sigma_p(s|\omega, \sigma_f)$ is continuous under pointwise weak convergence of σ_f .

To sum up, we have represented our game as a static game between four players (two parties, a representative voter, and a fictional player), who choose strategies

$$(x_a, x_b, \sigma, \sigma_f) \in (\mathbb{R}^n)^S \times (\mathbb{R}^n)^S \times (\Delta((\mathbb{R}^n)^S))^{\Omega \times D} \times (\Delta((\mathbb{R}^n)^S))^{\Omega \times D}.$$

Payoffs: The payoff U_a of party a is

$$U_a(x_a, x_b, \sigma, \sigma_f) := \sum_{\omega, \delta, s} \mu(\omega, \delta) \sigma_p(s|\omega, \sigma_f) \left(m \int F_{\nu} \left(u(x_a(s), p_S(s)) - u(x_b(s), p_S(s)) \right) d\sigma(p_S|\omega, \delta) + u(x_a(s), x_a^*) \right).$$

The payoff U_b for party b is defined analogously.

The utility U_v of the representative voter is

$$U_v(x_a, x_b, \sigma, \sigma_f) := \sum_{\omega, \delta, s} \mu(\omega, \delta) \sigma_p(s|\omega, \sigma_f) \int v(p_S(s), \omega + \delta, x_a(s), x_b(s)) d\sigma(p_S(s)|\omega, \delta) - c(\sigma) \quad (23)$$

where

$$\begin{aligned} v(p, \theta, x, y) &:= \int_{\mathbb{R}} \begin{cases} u(x, \theta) & \text{if } u(x, p) \geq u(y, p) + \nu \\ u(y, \theta) + \nu & \text{if } u(x, p) < u(y, p) + \nu \end{cases} dF(\nu) \\ &= F_{\nu}(u(x, p) - u(y, p))u(x, \theta) + (1 - F_{\nu}(u(x, p) - u(y, p)))u(y, \theta) \\ &\quad + \int_{u(x, p) - u(y, p)}^{\infty} \nu dF(\nu), \\ c(\sigma) &:= D_{\text{KL}}(P^{(\omega, \delta), p_S} || P^{(\omega, \delta)} \otimes P^{p_S}). \end{aligned}$$

The utility $v(p, \theta, x, y)$ captures that the voter votes optimally between platforms x and y given reported posterior mean p . The information cost $c(\sigma)$ is mutual information, which is the Kullback-Leibler divergence of the joint distribution $P^{(\omega, \delta), p_S}$ of state (ω, δ) and voter signal p_S from the product distribution $P^{(\omega, \delta)} \otimes P^{p_S} = \mu \otimes P^{p_S}$.

The fictional fourth player has payoff

$$U_f(x_a, x_b, \sigma, \sigma_f) := \begin{cases} 1 & \text{if } \sigma_f = \sigma, \\ 0 & \text{else.} \end{cases}$$

Thus, in equilibrium $\sigma_f = \sigma$.

Proof. First, we prove equilibrium existence through a fixed point theorem. Second, we prove existence of an equilibrium in which the desired statement of Theorem 4 holds.

Fixed Point Theorem: We apply the Kakutani-Fan-Glicksberg fixed point theorem (Aliprantis and Border, 2006, Theorem 17.55) to show existence of a pure-strategy equilibrium. It states that a correspondence Φ with closed graph and nonempty convex values on a nonempty compact convex subset K of a locally convex Hausdorff space has a fixed point.

Below, we define K as a nonempty compact convex subset of the strategy space

$$(\mathbb{R}^n)^S \times (\mathbb{R}^n)^S \times (\Delta((\mathbb{R}^n)^S))^{\Omega \times D} \times (\Delta((\mathbb{R}^n)^S))^{\Omega \times D},$$

by ruling out certain dominated strategies. Our strategy spaces are metrizable and hence Hausdorff. The weak topology is induced by a family of seminorms (the integral with respect to continuous bounded functions) and hence locally convex.

We construct Φ as the best-response correspondence. We show below that the best-response correspondences are upper hemicontinuous and nonempty compact-valued. Thus, by the closed graph theorem, the graph of Φ is closed. Finally, we show that the best-response correspondences are convex-valued, through showing that the payoffs are concave.

Compact and Convex Strategy Spaces: While strategy spaces are not compact, we can restrict attention to compact convex spaces of undominated strategies.

As shown in Appendix D.7, party j would never choose a platform outside a certain compact ellipse around their ideal points, $\mathcal{E}_j$. By Tychonoff's theorem, the strategy space $\mathcal{E}_j^S$ is compact (under the topology of pointwise convergence) and because $\mathcal{E}_j$ is compact. It is convex because $\mathcal{E}_j$ is convex.

For a voter it is never optimal to report a posterior mean that lies outside of the convex hull $\text{conv } \Theta$ of the support Θ because the posterior mean given any belief must lie inside this convex hull. Because Θ is finite, $\text{conv } \Theta$ is compact, and therefore $(\text{conv } \Theta)^S$ is compact. Because $(\text{conv } \Theta)^S$ is also metrizable, $\Delta((\text{conv } \Theta)^S)$ is compact under the topology of weak convergence. The set $(\Delta((\text{conv } \Theta)^S))^{\Omega \times D}$ is compact under the pointwise topology of weak convergence.

The fictional fourth player's best response is never outside the compact space of strategies $(\Delta((\text{conv } \Theta)^S))^{\Omega \times D}$ of the representative voter.

Formally, we can restrict attention to the compact and convex strategy space

$$K := \mathcal{E}_a^S \times \mathcal{E}_b^S \times (\Delta((\text{conv } \Theta)^S))^{\Omega \times D} \times (\Delta((\text{conv } \Theta)^S))^{\Omega \times D}.$$

Upper Hemicontinuous Best Response Correspondence: We show that the voter objective U_v is upper semicontinuous in $(x_a, x_b, \sigma, \sigma_f)$ and continuous in (x_a, x_b, σ_f) . Together with the feasibility set of σ being nonempty compact-valued and constant, it follows from the generalization of Berge's maximum theorem due to Tian and Zhou (1992), analogous to our proof of Proposition 7, that the best-response correspondence is nonempty compact-valued and upper hemicontinuous.

As a first step, the instrumental value of information is jointly continuous in players' strategies. The function v in (23) is uniformly continuous in (x_a, x_b) due to continuous differentiability over a compact domain, and continuous in p by continuity of f_ν . Moreover, v is bounded by the compact domain of (p, θ, x_a, x_b) and by ν having finite absolute first moment. Thus, by the Portmanteau theorem, the integral in (23) is continuous under weak convergence of $\sigma(p_S(s)|\omega, \delta)$. By uniform convergence of the integrand in $(x_a(s), x_b(s))$, the integral is jointly continuous in $(\sigma, x_a(s), x_b(s))$ (see (42)). By continuity of $\sigma_p(s|\omega, \sigma_f)$ in σ_f and because the sum over (ω, δ, s) in (23) is finite, the instrumental value is jointly continuous in players' strategies.

As a second step, the information cost, which depends only on σ , is lower semicontinuous, making the voter objective jointly upper semicontinuous. By Posner (1975), the Kullback-Leibler divergence $D_{\text{KL}}(P||Q)$ is jointly lower semicontinuous under weak convergence of P and Q . The joint distribution $P^{(\omega, \delta), p_S}$ is just a finite average of the conditional distributions of p_S conditional on (ω, δ) , that is, $\sigma(\cdot|\omega, \delta)$, so the joint distribution weakly converges as σ weakly converges pointwise. Similarly, the distribution P^{p_S} converges weakly as σ does. The product measure $P^{(\omega, \delta)} \otimes P^{p_S}$ converges weakly if P^{p_S} does, which can be verified via the Portmanteau theorem by testing expectations $\mathbb{E}_{(\omega, \delta), p_S}[f((\omega, \delta), p_S)]$ under $((\omega, \delta), p_S) \sim P^{(\omega, \delta)} \otimes P^{p_S}$ of continuous bounded functions f . Such expectations converge because they are weighted averages of expectations that converge by the Portmanteau theorem, $\mathbb{E}_{(\omega, \delta), p_S}[f((\omega, \delta), p_S)] = \sum_\theta P^{(\omega, \delta)}(\omega, \delta) \mathbb{E}_{p_S}[f((\omega, \delta), p_S)]$.

Combining the first and second step, the voter objective is upper semicontinuous in $(x_a, x_b, \sigma, \sigma_f)$ and continuous in (x_a, x_b, σ_f) .

The party objective U_j is jointly continuous in players' strategies by an analogous argument to the voter's instrumental value of information being jointly continuous. Thus, by Berge's maximum theorem, the best-response correspondence of parties is upper hemicontinuous and nonempty compact-valued and therefore has a closed graph.

The fictional player's utility is not continuous but their best response $\sigma_f = \sigma$ is nevertheless continuous in (x_a, x_b, σ) .

Convex-valued Best Response Correspondence: The set of best responses σ to (σ_f, x_a, x_b) is convex, because U_v is concave in σ . The instrumental value of information is linear in σ and the Kullback-Leibler divergence $c(\sigma)$ is convex in the conditional distribution σ : both the joint and the product measure are linear in the conditional distribution σ and the Kullback-Leibler divergence is convex.

For the party objective to be concave in x_a , we require again that m is small enough or the valence shock is large enough, see Appendix D.7. By compactness of $\text{conv } \Theta$, the set of distributions of posterior means supported on $\text{conv } \Theta$ is compact. Then, by an argument as in Lemma 27, for m small enough or valence ν large enough, the party objective is strictly concave and the best response correspondence is single- and therefore convex-valued.

This concludes our proof of equilibrium existence. Next, we show that there exists an equilibrium where party platforms respond to ω only through $(x_b^* - x_a^*)^\top A\omega$.

Existence of an equilibrium with the desired properties: We show that there exists an equilibrium

in which voters' acquired signal structures do not distinguish between aggregate states ω that have the same A -projection on $x_b^* - x_a^*$, $\langle x_b^* - x_a^*, \omega \rangle_A$. Because party platforms respond to voter preferences (Lemma 1), this implies that party platforms do not distinguish between such states. Formally, there exists an equilibrium $(x_a, x_b, \sigma, \sigma_f)$ such that σ satisfies the measurability condition

$$\forall \omega, \omega', \delta, p_S: \langle x_b^* - x_a^*, \omega \rangle_A = \langle x_b^* - x_a^*, \omega' \rangle_A \rightarrow \sigma(p_S | \omega, \delta) = \sigma(p_S | \omega', \delta). \quad (24)$$

To show this, we restrict σ and σ_f to signal structures that satisfy the measurability condition (24) and x_a and x_b to functions such that the platform difference is necessarily parallel to the ideological difference of parties,

$$\forall s: x_a(s) - x_b(s) \parallel x_a^* - x_b^*. \quad (25)$$

Formally, let $K' \subset K$ be the subset of strategies $(x_a, x_b, \sigma, \sigma_f) \in K$ that satisfy both the measurability condition (24) and the parallelity condition (25). The space K' is nonempty convex subset of K . Because K' is a closed subspace of the compact space K , K' is compact. Then, we show that the best-response correspondence restricted to the subspace K' maps into K' . By the Kakutani-Fan-Glicksberg fixed-point theorem, there is an equilibrium where (24) holds.

First, parties' best responses satisfies the parallelity condition (25) by Lemma 1. Because parties care about the expected vote share, Lemma 1 also holds for any *belief* over the distribution of voter preferences that parties share.

Second, the best response of σ_f satisfies the measurability condition (24) if σ does, because the best response is simply $\sigma_f = \sigma$.

Third and finally, we prove that any best response σ to (x_a, x_b, σ_f) satisfies (24) if σ_f satisfies (24) and (x_a, x_b) satisfy (25). To prove this, we show that aggregate states with the same projection on the ideological difference of parties are payoff equivalent. Because of the invariance property of mutual information, voters optimal learning strategy does not distinguish between payoff equivalent states.

Let ω and ω' be such that $\langle x_b^* - x_a^*, \omega \rangle_A = \langle x_b^* - x_a^*, \omega' \rangle_A$. If σ_f satisfies (24), then the public signal σ_p satisfies the measurability condition $\sigma_p(s | \omega, \sigma_f) = \sigma_p(s | \omega', \sigma_f)$ by

$$\begin{aligned} \sigma_p(s | \omega, \sigma_f) &= \sigma_p \left(s \middle| \sum_{\delta} \mu(\delta) \sigma_f(\cdot | \omega, \delta) \right) \\ &= \sigma_p \left(s \middle| \sum_{\delta} \mu(\delta) \sigma_f(\cdot | \omega', \delta) \right) = \sigma_p(s | \omega', \sigma_f). \end{aligned} \quad (26)$$

By $\langle x_b^* - x_a^*, \omega \rangle_A = \langle x_b^* - x_a^*, \omega' \rangle_A$, we have, for any δ , $\langle x_b^* - x_a^*, \theta \rangle_A = \langle x_b^* - x_a^*, \theta' \rangle_A$ where $\theta = \omega + \delta$

and $\theta' = \omega' + \delta$. By the parallelity condition (25), this implies for all $s \in S$

$$\begin{aligned} u(x_a(s), \theta) - u(x_b(s), \theta) &= \left\langle x_a(s) - x_b(s), \theta - \frac{x_a(s) + x_b(s)}{2} \right\rangle_A \\ &= \left\langle x_a(s) - x_b(s), \theta' - \frac{x_a(s) + x_b(s)}{2} \right\rangle_A = u(x_a(s), \theta') - u(x_b(s), \theta'). \end{aligned}$$

Therefore,

$$\begin{aligned} v(p, \theta, x_a(s), x_b(s)) &= F_\nu(u(x_a(s), p) - u(x_b(s), p)) \underbrace{(u(x_a(s), \theta) - u(x_b(s), \theta))}_{= u(x_a(s), \theta') - u(x_b(s), \theta')} + u(x_b(s), \theta) \\ &\quad + \int_{u(x_a(s), p) - u(x_b(s), p)}^\infty \nu dF(\nu) \\ &= v(p, \theta', x_a(s), x_b(s)) + u(x_b(s), \theta) - u(x_b(s), \theta'). \end{aligned} \tag{27}$$

The voter's utility, neglecting the information cost, under action p_S and state (ω, δ) , is

$$\sum_s \sigma_p(s|\omega, \sigma_f) v(p_S(s), \omega + \delta, x_a(s), x_b(s)).$$

Equations (26) and (27) imply that this voter's utility is the same under state (ω', δ) —up to a constant that does not interact with the action and is therefore immaterial:

$$\begin{aligned} &\sum_s \sigma_p(s|\omega, \sigma_f) v(p_S(s), \omega + \delta, x_a(s), x_b(s)) \\ &= \sum_s \sigma_p(s|\omega, \sigma_f) \left(v(p_S(s), \omega' + \delta, x_a(s), x_b(s)) - u(x_b(s), \theta') + u(x_b(s), \theta) \right) \\ &= \sum_s \sigma_p(s|\omega', \sigma_f) \left(v(p_S(s), \omega' + \delta, x_a(s), x_b(s)) \right) + \sum_s \sigma_p(s|\omega', \sigma_f) \left(u(x_b(s), \theta) - u(x_b(s), \theta') \right) \end{aligned}$$

This shows that states (ω, δ) and (ω', δ) are payoff-equivalent. By information monotonicity of mutual information (Amari, 2016; Caplin, Dean, and Leahy, 2022), any optimal signal structure σ does not distinguish between these states in the sense of (24). $\square$

B Appendix: Alternative Timing

B.1 One-Dimensional Policy Space

B.1.1 Results on Voter Learning

This section characterizes the optimal voter learning strategies and resulting expected vote shares given party platforms $x_a, x_b \in \mathbb{R}$.

As discussed in the main text, the voter learning problem can be expressed as a function of the distribution $\rho \in \Delta(\mathbb{R})$ over posterior means. As is known, such a distribution ρ can be generated by a Bayes-consistent distribution $\tau \in \Delta(\Delta(\mathbb{R}))$ over posteriors if and only if $\rho \leq_{\text{MPS}} \mu$. The following definition will be useful for the following.

Definition 1 (Feasibility). *We say a binary set $\{\theta_a, \theta_b\}$ is **feasible** if there exists a distribution ρ over $\{\theta_a, \theta_b\}$, $\rho \in \Delta(\{\theta_a, \theta_b\})$, that is a mean-preserving contraction of the prior μ , $\rho \leq_{\text{MPS}} \mu$.*

Note that if ρ is a mean-preserving contraction of μ , then ρ has the same expectation as μ , which is 0. Thus, if $0 \notin [\theta_a, \theta_b]$, then $\{\theta_a, \theta_b\}$ is not feasible. On the other hand, if $0 \in [\theta_a, \theta_b]$, then there is a unique distribution ρ over $\{\theta_a, \theta_b\}$ with expectation zero. Thus, feasibility boils down to whether this distribution ρ is a mean-preserving contraction of the prior μ .

Without loss, assume $x_a \leq x_b$ and define

$$\begin{aligned}\theta_a(x_a, x_b) &:= \frac{x_a + x_b}{2} - \frac{x_b - x_a}{2\kappa}, \\ \theta_b(x_a, x_b) &:= \frac{x_a + x_b}{2} + \frac{x_b - x_a}{2\kappa}.\end{aligned}$$

Definition 2 (Binding feasibility). *We say **feasibility is not binding at** $(\mathbf{x}_a, \mathbf{x}_b)$ if the posterior means $\{\theta_a(x_a, x_b), \theta_b(x_a, x_b)\}$ are feasible. Otherwise, we say feasibility is binding at (x_a, x_b) .*

The following proposition characterizes the solution to the voter learning problem.

Proposition 4. *A solution $\rho \in \Delta(\mathbb{R})$ to the voter learning problem given $x_a, x_b \in \mathbb{R}$ exists and is unique. Let $\theta = \theta(x_a, x_b)$ and $\theta_b = \theta_b(x_a, x_b)$. There are three cases:*

- (1) *If $\{\theta_a, \theta_b\}$ is feasible, voters acquire the two posterior means θ_a and θ_b .*
- (2) *If $\theta_a < 0 < \theta_b$ and $\{\theta_a, \theta_b\}$ is not feasible, voters acquire a threshold signal structure $s(\theta) = \mathbb{1}(\theta > t)$, where $t \in \mathbb{R} \cup \{-\infty, \infty\}$ is more extreme than the party midpoint, that is, we have $0 < \frac{x_a + x_b}{2} < t$, $0 = \frac{x_a + x_b}{2} = t$, or $t < \frac{x_a + x_b}{2} < 0$. As κ decreases, the threshold t moves weakly closer to the party midpoint $\frac{x_a + x_b}{2}$.*
- (3) *If $0 \notin (\theta_a, \theta_b)$, voters acquire no information.*

When voters acquire information, then one posterior mean θ_a is always closer to x_a and the other, θ_b , is closer to x_b . The expected vote shares of parties are simply the probabilities of posterior means θ_a and θ_b . These probabilities are uniquely pinned down by $\Pr_\rho(\theta_a)\theta_a + \Pr_\rho(\theta_b)\theta_b = 0$.

The following result summarizes the implications for parties' expected vote shares and is crucial for the analysis of endogenous party positions.

Corollary 3. *There are three cases.*

(1) *If $\{\theta_a, \theta_b\}$ is feasible, the expected vote share P_a of party a is*

$$P_a(x_a, x_b) = \frac{1}{2} + \frac{\kappa x_b + x_a}{2 x_b - x_a}.$$

(2) *If $\theta_a < 0 < \theta_b$ and $\{\theta_a, \theta_b\}$ is not feasible, the expected vote share P_a of party a is $P_a = F_\mu(t)$, where t is as in Proposition 4. If party a is more extreme, that is, $|x_a| > |x_b|$, then $P_a(x_a, x_b) < \frac{1}{2} + \frac{\kappa x_b + x_a}{2 x_b - x_a}$.*

(3) *If $0 \notin (\theta_a, \theta_b)$, the party whose position is closer to 0 obtains all votes.*

Interestingly, the expected vote share in the first case takes the same form as in the moderate utility model (Halff, 1976; He and Natenzon, 2024).

Finally, the following corollary shows that costly voter learning generates a “bias” toward the moderate party: the moderate party receives more votes than the true share of voters whose ideal points are closer to its position.

Corollary 4 (Bias toward Moderate Party). *Under $|x_a| < |x_b|$, party a obtains a higher expected vote share P_a than the true share of voters that are closer to x_a than to x_b , that is, $P_a > F_\mu(\frac{x_a + x_b}{2})$. The expected vote share of party a decreases as the information cost parameter κ decreases or as party polarization $|x_b - x_a|$ increases, holding $\frac{x_a + x_b}{2}$ constant.*

B.1.2 Proposition 4

Proof. Using (16), the voter's maximization problem can be written, up to a constant, as a function of the acquired distribution ρ over posterior means. By Strassen's theorem, ρ is feasible if it is a mean-preserving contraction of the prior μ . The voter solves

$$\begin{aligned} \max_{\rho \in \Delta(\mathbb{R})} \mathbb{E}_{\theta \sim \rho} \left[\max \left\{ (x_b - x_a) \left(\theta - \frac{x_a + x_b}{2} \right), -(x_b - x_a) \left(\theta - \frac{x_a + x_b}{2} \right) \right\} - \kappa \theta^2 \right] \\ \text{s.t. } \rho \leq_{\text{MPS}} \mu \end{aligned}$$

We first solve the relaxed problem where we relax the constraint $\rho \leq_{\text{MPS}} \mu$ to $\mathbb{E}_{\theta \sim \rho}[\theta] = 0$. Through the coordinate change

$$\hat{\theta} = \theta - \frac{x_a + x_b}{2},$$

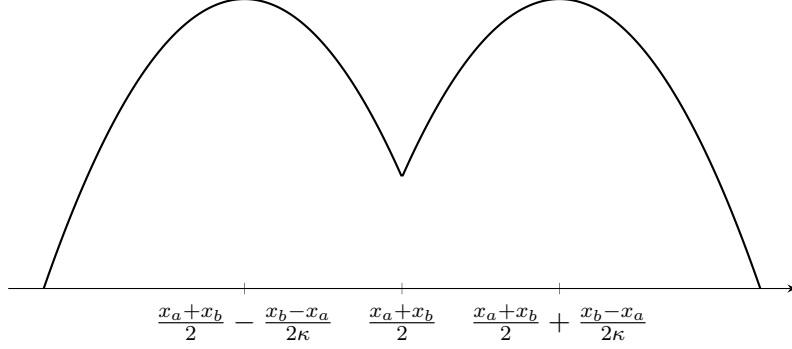


Figure 6: The value function $v(\theta)$

the voter's problem becomes

$$\begin{aligned} \max_{\rho \in \Delta(\mathbb{R})} \mathbb{E}_{\hat{\theta} \sim \hat{\rho}} \left[\overbrace{\max \left\{ (x_b - x_a)\hat{\theta}, -(x_b - x_a)\hat{\theta} \right\}}^{v(\hat{\theta})} - \kappa \hat{\theta}^2 \right] \\ \text{s.t. } \mathbb{E}_{\hat{\theta} \sim \hat{\rho}}[\theta] = -\frac{x_a + x_b}{2} \end{aligned}$$

Because the value function v is symmetric around 0 (see Figure 6), the solution is obtained simply by finding its maxima. The gradient of the value function is zero if $2\kappa\hat{\theta} = x_b - x_a$ or

$$\hat{\theta} = \frac{x_b - x_a}{2\kappa}.$$

Define

$$\begin{aligned} \theta_a &:= \frac{x_a + x_b}{2} - \frac{x_b - x_a}{2\kappa}, \\ \theta_b &:= \frac{x_a + x_b}{2} + \frac{x_b - x_a}{2\kappa}. \end{aligned}$$

If the origin lies between θ_a and θ_b , then the voter acquires the posterior means θ_a and θ_b .

This is the solution if $\{\theta_a, \theta_b\}$ is feasible. If $\{\theta_a, \theta_b\}$ is not feasible, then a threshold signal $s(\theta) = \mathbb{1}(\theta > t)$ is optimal by the following. The distribution ρ over $\{\theta_a, \theta_b\}$ is the only solution to the voter learning problem which satisfies the first-order condition of optimality of the relaxed problem (see the Lagrangian Lemma in Caplin, Dean, and Leahy, 2022). Thus, the solution to the non-relaxed problem must have a binding mean-preserving contraction constraint. The mean-preserving contraction constraint $\rho \leq_{\text{MPS}} \mu$ can, by Strassen's theorem, be written as

$$\forall t \in \mathbb{R} : \int_{-\infty}^t F_{\rho}(\theta) d\theta \leq \int_{-\infty}^t F_{\mu}(\theta) d\theta.$$

If ρ is binary and the constraint is binding at some t , then ρ is induced by the threshold signal with threshold t . Also, any threshold signal with threshold t induces a distribution ρ with binding mean-preserving contraction constraint. Note that we allow the threshold signal to have threshold $t = \pm\infty$, in which case the voter acquires no information and unconditionally votes for the party closer to the origin.

If a threshold signal is optimal, the threshold t is more extreme than $\frac{x_a+x_b}{2}$ by the following argument. It is easy to see that the *instrumental value of information* $b(t)$ is increasing as t moves closer to $\frac{x_a+x_b}{2}$:

$$\begin{aligned} b(t) &:= \int_{-\infty}^t -(x_a - \theta)^2 f(\theta) d\theta + \int_t^{\infty} -(x_b - \theta)^2 f(\theta) d\theta \\ \Rightarrow b'(t) &= (-(x_a - t)^2 + (x_b - t)^2) f(t) \end{aligned}$$

Note that $b'(\frac{x_a+x_b}{2}) = 0$.

The *information cost* decreases as t becomes more extreme if the prior is log-concave. The information cost is

$$\begin{aligned} c(t) &:= \text{Var}[\theta] - (F(t) \text{Var}[\theta|\theta \leq t] + (1 - F(t)) \text{Var}[\theta|\theta > t]) \\ &= \text{Var}[\theta] - \left(F(t) \min_E \left\{ \frac{\int_{-\infty}^t (\theta - E)^2 dF(\theta)}{F(t)} \right\} + (1 - F(t)) \min_E \left\{ \frac{\int_t^{\infty} (\theta - E)^2 dF(\theta)}{1 - F(t)} \right\} \right) \\ &= \text{Var}[\theta] - \left(\min_E \left\{ \int_{-\infty}^t (\theta - E)^2 f(\theta) d\theta \right\} + \min_E \left\{ \int_t^{\infty} (\theta - E)^2 f(\theta) d\theta \right\} \right) \end{aligned}$$

Using the envelope theorem, we obtain the derivative

$$\begin{aligned} c'(t) &= f(t)(t - \mathbb{E}[\theta|\theta > t])^2 - f(t)(t - \mathbb{E}[\theta|\theta \leq t])^2 \\ &= f(t)(\mathbb{E}[\theta - t|\theta > t]^2 - \mathbb{E}[t - \theta|\theta \leq t]^2). \end{aligned}$$

By Theorem 1.C.52 in Shaked and Shanthikumar (2007), the conditional expectation $\mathbb{E}[\theta - t|\theta > t]$ is decreasing in t for logconcave density of θ . By symmetry of f around 0, $\mathbb{E}[t - \theta|\theta \leq t] = \mathbb{E}[\theta - (-t)|\theta > -t]$. Together, this delivers for $t > 0$,

$$0 < \mathbb{E}[\theta - t|\theta > t] < \mathbb{E}[\theta|\theta > 0] < \mathbb{E}[\theta - (-t)|\theta > (-t)] = \mathbb{E}[t - \theta|\theta \leq t].$$

Thus, $c'(t)$ is negative if $t > 0$.

Note that for $\frac{x_a+x_b}{2} > 0$, $b'(t) > 0$ and $c'(t) < 0$ for all $t < \frac{x_a+x_b}{2}$, so $t = \frac{x_a+x_b}{2}$ dominates any $t < \frac{x_a+x_b}{2}$. Moreover, $b'(\frac{x_a+x_b}{2}) = 0$ and $c'(\frac{x_a+x_b}{2}) < 0$, so $t > \frac{x_a+x_b}{2}$ is optimal. Recall that the optimal distribution over posteriors is unique and thus the optimal t is unique.

Restrict t to $[\frac{x_a+x_b}{2}, \infty]$ ($t = \infty$ refers to always voting for party a). The utility $b(t) - \kappa c(t)$ is supermodular in (κ, t) by $c'(t) < 0$. Thus, the optimal t is weakly increasing as a function of κ . $\square$

B.1.3 Corollary 3

Proof. The case (3) is immediate.

Case (1): By the law of iterated expectations, we have

$$P_a \left(\frac{x_a + x_b}{2} - \frac{x_b - x_a}{2\kappa} \right) + (1 - P_a) \left(\frac{x_a + x_b}{2} + \frac{x_b - x_a}{2\kappa} \right) = \mathbb{E}[\theta] = 0 \Rightarrow P_a = \frac{1}{2} + \frac{\kappa}{2} \frac{x_a + x_b}{x_b - x_a}.$$

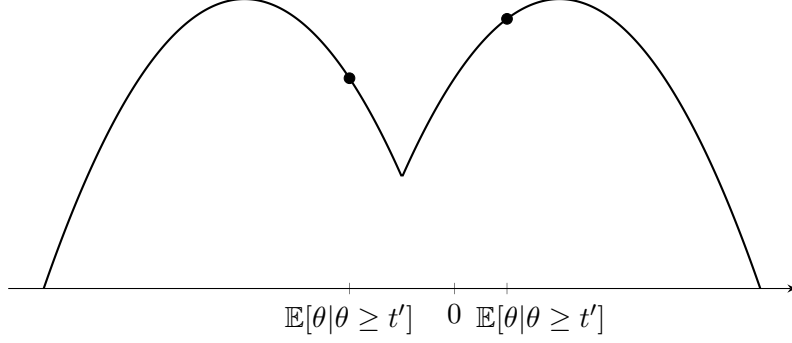


Figure 7: Value function and feasible posterior means under threshold signal t'

Case (2): That $P_a = F_\mu(t)$ is immediate from case (2) of Proposition 1. The second part of the statement follows from the following lemma.

Lemma 6. *Let x_a, x_b be such that $\{\theta_a, \theta_b\}$ is not feasible, $|x_a| > |x_b|$, and $x_a < x_b$. The expected vote share P_a of party a satisfies $P_a < \frac{1}{2} + \frac{\kappa}{2} \frac{x_a + x_b}{x_b - x_a}$.*

Proof. First, t is such that the induced posterior means $\mathbb{E}[\theta|\theta < t]$ and $\mathbb{E}[\theta|\theta > t]$ are in between the maxima $\{\theta_a, \theta_b\}$ of the value function,

$$\theta_a \leq \mathbb{E}[\theta|\theta < t] < \mathbb{E}[\theta|\theta > t] \leq \theta_b, \quad (28)$$

or the voter acquires no information. The reason is the following. If the posterior means were outside the maxima of the value function, then the maxima would be a garbling of the threshold signal, so the maxima would be feasible. Both posterior means need to be on different sides of $\frac{x_a + x_b}{2}$, otherwise the agent takes the same action regardless of the information they acquired and no information would be better. Finally, it cannot be the case that

$$\mathbb{E}[\theta|\theta < t] < \theta_a \leq \frac{x_a + x_b}{2} \leq \mathbb{E}[\theta|\theta > t] \leq \theta_b,$$

because increasing the threshold t marginally would increase both posterior means $\mathbb{E}[\theta|\theta < t]$ and $\mathbb{E}[\theta|\theta > t]$, which would result in a higher value because the derivative of the value function is positive at $\mathbb{E}[\theta|\theta < t]$ and non-negative at $\mathbb{E}[\theta|\theta > t]$, see also Figure 7. Analogously, the posterior mean $\mathbb{E}[\theta|\theta > t]$ cannot be to the right of θ_b . Thus, the posterior means must satisfy the ordering (28) or $t = \pm\infty$.

Second, consider the threshold t' that delivers the same vote shares as $\tilde{P}_a = \frac{1}{2} + \frac{\kappa}{2} \frac{x_a + x_b}{x_b - x_a}$, $t' = F_\mu^{-1}(P_a)$. By $|x_a| > |x_b|$, $\tilde{P}_a < 1/2$, so by symmetry of μ , $t' < 0$. By the law of iterated expectations

$$0 = \tilde{P}_a \mathbb{E}[\theta|\theta < t'] + (1 - \tilde{P}_a) \mathbb{E}[\theta|\theta \geq t'] = \tilde{P}_a \theta_a + (1 - \tilde{P}_a) \theta_b.$$

Hence, there exists an $\alpha > 0$ such that the induced posterior means satisfy

$$(\mathbb{E}[\theta|\theta < t'], \mathbb{E}[\theta|\theta \geq t']) = (\alpha \theta_a, \alpha \theta_b).$$

By $\{\theta_a, \theta_b\}$ not being feasible, $\alpha < 1$. We show that the optimal t satisfies $t < t'$, that is, the binding mean-preserving contraction constraint benefits the moderate party. For that we use the following Lemma.

Lemma 7 (Jewitt's Lemma). $\Delta(t) := \mathbb{E}[\theta|\theta > t] - \mathbb{E}[\theta|\theta < t]$ is weakly decreasing in t for $t < 0$ and weakly increasing in t for $t > 0$.

Proof. Jewitt (2004) shows that $\Delta(t)$ is quasiconvex if the density of θ is quasiconcave. Because θ has a symmetric density around 0, $\Delta(t)$ is weakly decreasing for $t < 0$ and weakly increasing for $t > 0$. $\square$

We prove $t < t'$ by showing that increasing t beyond t' lowers the voter's utility. Intuitively, increasing t has two effects on the voter's utility: it shifts the posterior means and it shifts their probabilities. We show that both effects have a negative effect on the voter's utility (consulting Figure 7 may help visualize the following steps). First, note that we have for $t > t'$,

$$\frac{-v'(\mathbb{E}[\theta|\theta < t])}{v'(\mathbb{E}[\theta|\theta \geq t])} \geq \frac{-v'(\mathbb{E}[\theta|\theta < t'])}{v'(\mathbb{E}[\theta|\theta \geq t'])} = \frac{1 - F(t')}{F(t')}.$$

A $t > 0$ is clearly suboptimal because $-t$ would deliver an equal cost of information but higher higher instrumental value of information. Under $t < 0$, Jewitt's Lemma implies

$$\frac{d\Delta(t)}{dt} < 0 \Rightarrow \frac{d\mathbb{E}[\theta|\theta < t]}{dt} > \frac{d\mathbb{E}[\theta|\theta \geq t]}{dt}.$$

Now, for any t with $t' \leq t \leq 0$, as long as the ordering requirement (28) holds,

$$\begin{aligned} & \frac{d}{dt} \left(F(t)v(\mathbb{E}[\theta|\theta < t]) + (1 - F(t))v(\mathbb{E}[\theta|\theta \geq t]) \right) \\ &= F(t)v'(\mathbb{E}[\theta|\theta < t])\frac{d\mathbb{E}[\theta|\theta < t]}{dt} + (1 - F(t))v'(\mathbb{E}[\theta|\theta \geq t])\frac{d\mathbb{E}[\theta|\theta \geq t]}{dt} + f(t)\underbrace{(v(\mathbb{E}[\theta|\theta < t]) - v(\mathbb{E}[\theta|\theta \geq t]))}_{<0} \\ &< \left(F(t)v'(\mathbb{E}[\theta|\theta < t]) + (1 - F(t))v'(\mathbb{E}[\theta|\theta \geq t]) \right) \frac{d\mathbb{E}[\theta|\theta \geq t]}{dt} + F(t)\underbrace{v'(\mathbb{E}[\theta|\theta < t])}_{<0} \underbrace{\left(\frac{d\mathbb{E}[\theta|\theta < t]}{dt} - \frac{d\mathbb{E}[\theta|\theta \geq t]}{dt} \right)}_{>0} \\ &< \underbrace{\left(F(t)v'(\mathbb{E}[\theta|\theta < t]) + (1 - F(t))v'(\mathbb{E}[\theta|\theta \geq t]) \right)}_{\leq 0} \underbrace{\frac{d\mathbb{E}[\theta|\theta \geq t]}{dt}}_{>0} \leq 0 \end{aligned}$$

Therefore, $t < t'$ is optimal. $\square$

This concludes the proof of Corollary 5. $\square$

B.1.4 Corollary 4

Proof. We go through the three cases of Corollary 3 by increasing size of κ . In each case, the expected vote share is biased toward the moderate party and this bias increases in κ . Increasing

party polarization $|x_b - x_a|$ while holding $\frac{x_a + x_b}{2}$ constant, has the same effect on the voter's objective as decreasing κ (and, at the same time, multiplying the objective by an irrelevant constant).

By Corollary 3, if (θ_a, θ_b) are not feasible, then the optimal signal is a threshold signal with a threshold more extreme than the party midpoint, thus there is a bias toward the moderate party. The threshold becomes more extreme as κ increases, so the bias increases.

Let κ be such that $\{\theta_a, \theta_b\} = \{\frac{x_a + x_b}{2} - \frac{\kappa}{2} \frac{x_a + x_b}{x_b - x_a}, \frac{x_a + x_b}{2} + \frac{\kappa}{2} \frac{x_a + x_b}{x_b - x_a}\}$, is just feasible. Then, these posterior means are induced by a threshold signal, which by the above, is biased toward the moderate party. For a higher κ , the posterior means $\{\theta_a, \theta_b\}$ move symmetrically closer to the party midpoint $\frac{x_a + x_b}{2}$. This increases the weight on the more likely posterior mean, so it increases the bias.

If κ is large enough, then voters acquire no information and just vote for the ex-ante more attractive party, which is the moderate party. $\square$

B.1.5 Lemmas on Party Positions

The following lemma describes properties of any equilibrium and partially characterize parties' best response. We use these lemmas in the proofs of Proposition 5 and Proposition 6 below.

Lemma 8. *If $x_a^* < x_b^*$, then $x_a \leq x_b$ in equilibrium. If $0 \leq x_b^*$, then $x_a \leq x_b^*$ in equilibrium. If $x_a^* \leq 0$, then $x_a^* \leq x_b$ in equilibrium.*

Proof. Recall that we restrict attention to equilibria where both parties obtain positive expected vote shares and thereby positive probabilities of implementing their policies. The proof relies on the fact that it is suboptimal for party $j \in \{a, b\}$ to choose a position x_j that is further away from x_j^* than the other parties' position x_{-j} is from x_j^* . The reason is, as argued above, that parties are policy motivated.

First, we show $x_a \leq x_b$. Suppose, for the sake of contradiction, that $x_a > x_b$ and $x_a^* < x_b^*$. If x_b is at least as close to x_b^* than x_a is, then, given $x_a^* < x_b^*$ and $x_b < x_a$, x_b is closer to x_a^* than is x_a , leading to a contradiction.

Second, we show $x_a \leq x_b^*$ if $0 \leq x_b^*$. Suppose, for the sake of contradiction, that $x_a > x_b^*$. For x_b to be at least as close to x_b^* than x_a is, we need $x_b \leq x_a$. If $x_b < x_a$, this contradicts $x_a < x_b$. The remaining possibility is $x_a = x_b$. By $0 \leq x_b^* < x_a = x_b$, this is not an equilibrium: party a could choose $x_b - \varepsilon$ for ε small enough and obtain the full vote share and a better policy.

Analogously, one can show that $x_a^* \leq x_b$ if $x_a^* \leq 0$. $\square$

Lemma 9. *If $x_a^* \leq 0$, then $2x_a^* \leq x_a$ in equilibrium. If $0 \leq x_b^*$, then $x_b \leq 2x_b^*$ in equilibrium.*

Proof. We prove the first statement by distinguishing between different cases for the position x_b . The proof of the second statement is analogous.

Consider first $x_b \leq 0$. As shown by Lemma 8, in equilibrium $x_b \geq x_a^*$. In equilibrium, x_a is at least as close to x_a^* as is x_b , so $2x_a^* \leq x_a$.

If $0 < x_b$ and $x_a < x_a^*$ and feasibility is not binding at (x_a, x_b) , then $P_a(x_a, x_b) = \frac{1}{2} + \frac{\kappa}{2} \frac{x_b + x_a}{x_b - x_a}$ is increasing in x_a (this can be easily seen by differentiating $P_a(x_a, x_b)$ in x_a). Thus, increasing x_a would increase both the vote share and the policy utility when winning, so x_a is suboptimal. Thus, in equilibrium, $2x_a^* \leq x_a^* \leq x_a$.

If $0 < x_b$ and $x_a < 2x_a^*$ is more extreme than x_b , that is, $|x_a| \geq |x_b^*|$, then $P_a(x_a, x_b) \leq \frac{1}{2}$. Then, $x_a = 0$ delivers a greater expected vote share for party a (namely a vote share greater than $1/2$) and a greater policy utility when winning, so x_a is suboptimal.

Consider, finally, that case that $x_b > 0$, feasibility is binding at $x_a < x_a^*$, and x_a is less extreme than x_b , that is, $|x_a| < |x_b|$. If x_a is less extreme than x_b , then either party a obtains expected vote share 1 (which we have ruled out) or a threshold signal is optimal with $t > \frac{x_a + x_b}{2} > x_a$ (that is, the signal is biased toward the moderate party a) by Proposition 4. One can easily show that the voter utility $b(t) - \kappa c(t)$ is supermodular in x_a and t , when t is restricted to $t > x_a$, so t increases in x_a . Again, both the expected vote share and the policy utility when winning increase when raising x_a , so $x_a < x_a^*$ is suboptimal. $\square$

Lemma 10. *If $x_a^* \leq 0 \leq x_b^*$ and $\kappa \geq 1$, then the unique equilibrium is $(x_a, x_b) = (0, 0)$.*

Proof. Under $\kappa \geq 1$, the implemented policy is zero if one party chooses policy 0. This follows from Corollary 3 and the fact that $\theta_a(x_a, x_b)$ and $\theta_b(x_a, x_b)$ are weakly between x_a and x_b for $\kappa \geq 1$. Therefore, $(x_a, x_b) = (0, 0)$ is an equilibrium.

By Lemma 8, $x_a \leq x_b$ in any equilibrium.

Suppose (x_a, x_b) is an equilibrium. The average implemented policy is weakly positive or weakly negative. Suppose without loss that it is weakly positive. Then party a can obtain policy 0 for certain by choosing $x_a = 0$. This weakly improves the utility of party a because the policy outcome is weakly better in expectation and has no variance. Thus, if the policy outcome under (x_a, x_b) is not 0 for certain, then this is a profitable deviation. Finally, if the policy outcome is 0 for certain, then at least one party must choose position 0. The other party must also choose position 0, otherwise only one party would obtain a positive vote share, violating our equilibrium definition. $\square$

Lemma 11. *If $0 < x_a^*$ and $0 < x_b^*$, then either $x_a = x_b$ or*

$$x_a^* \leq x_a \leq x_b \leq 2x_b^* - x_a^*.$$

In particular, $|x_b - x_a| \leq 2|x_b^ - x_a^*|$.*

Proof. By Lemma 8, $x_a \leq x_b$ and by Lemma 9, $x_b \leq 2x_b^*$.

We show that $x_a = x_b$ or $x_a^* \leq x_a$ in an equilibrium with positive expected vote shares. If $x_a < x_a^*$ and $x_a < x_b$, then (x_a, x_b) is not an equilibrium because $U_a(x_a, x_b)$ could be improved. If x_a is more extreme than x_b , that is, $|x_a| \geq |x_b|$, then x_a must be negative and $P_a < 1/2$. Then, $x_a' = 0$ delivers higher policy utility when winning and a higher vote share of at least $1/2$, so x_a is suboptimal. If x_a is less extreme than x_b , that is, $|x_a| < |x_b|$, then increasing x_a slightly improves the policy utility when winning. Increasing x_a slightly also increases the vote share. As argued

in the proof of Lemma 26, if x_a is less extreme than x_b and $x_a < x_b$, then either party a obtains expected vote share 1 (which we have ruled out) or a threshold signal is optimal with threshold $t > \frac{x_a + x_b}{2} > x_a$ (that is, the signal is biased toward the moderate party a) by Proposition 4. One can easily show that the voter utility $b(t) - \kappa c(t)$ is supermodular in x_a and t , when t is restricted to $t > x_a$, so t increases in x_a .

Thus, if $x_a \neq x_b$, then $x_b < x_b^* + (x_b^* - x_a^*)$. By the above, if $x_a \neq x_b$, then $x_a^* \leq x_a$. By Lemma 8, $x_a \leq x_b^*$ in equilibrium, so together $x_a^* \leq x_a \leq x_b^*$. In an equilibrium with positive expected vote shares, x_b must be at least as close to x_b^* as is x_a . Because x_a has at most distance $|x_b^* - x_a^*|$ from x_b^* , x_b has at most distance $|x_b^* - x_a^*|$ from x_b^* in equilibrium, so $x_b < 2x_b^* - x_a^*$. $\square$

Next, we provide a lemma that partially characterizes parties best responses. We need a few definitions. Define

$$\begin{aligned}\tilde{P}_a(x_a, x_b) &:= \frac{1}{2} + \frac{\kappa}{2} \frac{x_b + x_a}{x_b - x_a}, \\ \tilde{U}_a(x_a, x_b) &:= \tilde{P}_a(x_a, x_b)u(x_a, x_a^*) + (1 - \tilde{P}_a(x_a, x_b))u(x_b, x_a^*).\end{aligned}$$

We call $\tilde{P}_a$ and $\tilde{U}_a$ the **pseudo vote share** and **pseudo utility**, respectively, because they are equal to the expected vote share and expected utility of party a only if feasibility is not binding at (x_a, x_b) , by Corollary 3.

Lemma 12 (Best Response). *Let $x_a^* < x_b^*$, $\kappa < 1$, and $x_b \in \mathbb{R}$. Define*

$$\hat{x}_a = x_a^* + \frac{\kappa}{1 - \kappa} x_b.$$

Necessity: *If x_a is the best response of party a to x_b and feasibility is not binding at (x_a, x_b) , then $x_a = \hat{x}_a$.*

Sufficiency: *If feasibility is not binding at $(\hat{x}_a, x_b)$ and $|\hat{x}_a| \geq |x_b|$, then $\hat{x}_a$ is the unique best response of party a to x_b .*

Analogous statements hold when exchanging a and b (maintaining $x_a^ < x_b^*$).*

Proof. Necessity: If feasibility is not binding at (x_a, x_b) , then party a 's objective is

$$\begin{aligned}\tilde{U}_a(x_a, x_b) &= \tilde{P}_a(x_a, x_b)u(x_a, x_a^*) + (1 - \tilde{P}_a(x_a, x_b))u(x_b, x_a^*) \\ &= (u(x_a, x_a^*) - u(x_b, x_a^*))\tilde{P}_a(x_a, x_b) + u(x_b, x_a^*)\end{aligned}\tag{29}$$

$$\begin{aligned}&= 2 \left(x_a^* - \frac{x_a + x_b}{2} \right) (x_a - x_b) \left(\frac{1}{2} + \frac{\kappa}{2} \frac{x_a + x_b}{x_b - x_a} \right) + u(x_b, x_a^*) \\ &= \left(\frac{x_a + x_b}{2} - x_a^* \right) (x_b(1 + \kappa) - x_a(1 - \kappa)) + u(x_b, x_a^*).\end{aligned}\tag{30}$$

The last term is a constant and can be ignored. Since the domain of x_a is unbounded and the objective any x_a further from x_a^* than x_b is dominated, the best response x_a to x_b must satisfy the first-order condition $\frac{dU_a(x_a, x_b)}{dx_a} = 0$. Simple derivation shows that the unique solution to the

first-order condition is

$$\hat{x}_a = x_a^* + \frac{\kappa}{1 - \kappa} x_b.$$

Sufficiency: To show that $\hat{x}_a$ is the best response if feasibility is not binding at $(\hat{x}_a, x_b)$ and $|\hat{x}_a| \geq |x_b|$, we rule out step-by-step that any other x_a could deliver a higher utility $U_a(x_a, x_b)$ than $\hat{x}_a$.

First, we rule out x_a such that feasibility is not binding at (x_a, x_b) . If feasibility is not binding at (x_a, x_b) , then party a 's objective is given by $\tilde{U}_a(x_a, x_b)$, which is negative quadratic in x_a and maximized at $\hat{x}_a$. Thus, $\hat{x}_a$ is the best response among all such x_a .

Inserting $\hat{x}_a$ into (30), we obtain that the resulting utility is

$$U_a(\hat{x}_a, x_b) = \frac{1 - \kappa}{2} \left(\frac{x_b}{1 - \kappa} - x_a^* \right)^2 + u(x_b, x_a^*) > u(x_b, x_a^*), \quad (31)$$

which we use below.

Second, we rule out x_a that are at least as far away from x_a^* than is x_b . Such x_a deliver utility $U_a(x_a, x_b) \leq u(x_b, x_a^*) < U_a(\hat{x}_a, x_b)$, using (31). Therefore, x_a delivers a lower utility than $\hat{x}_a$. Intuitively, because party a cares about the implemented policy, it never makes sense to propose a policy that is worse than its opponent's policy.

Third, we rule out x_a that deliver an expected vote share $P_a(x_a, x_b)$ of 0 or 1. If $P_a(x_a, x_b) = 0$, then the resulting utility is $U_a(x_a, x_b) = u(x_b, x_a^*) < U_a(\hat{x}_a, x_b)$, using (31). Therefore, x_a delivers a lower utility than $\hat{x}_a$. Consider $P_a(x_a, x_b) = 1$, which can be the case if 0 is not between $\theta_a(x_a, x_b)$ and $\theta_b(x_a, x_b)$ or if $\{\theta_a(x_a, x_b), \theta_b(x_a, x_b)\}$ is not feasible, by Corollary 3. We show that $1 < \tilde{P}_a(x_a, x_b)$. Assume, for the sake of contradiction, that x_a delivers a higher utility than $\hat{x}_a$. x_a can only deliver an expected vote share of 1 if $|x_a| < |x_b|$. By assumption we have $|x_b| < |\hat{x}_a|$. Because $\hat{x}_a$ and x_a deliver greater utility U_a than $u(x_b, x_a^*)$, $\hat{x}_a$ and x_a are closer to x_a^* than is x_b , by (31). This leaves only the case that $\hat{x}_a$ and x_b are on different sides of the origin and x_a is between $\hat{x}_a$ and x_b . This implies that if x_a delivers an expected vote share of 1, then $\tilde{P}_a > 1$. The reason is that if 0 is between θ_a and θ_b , then θ_a and θ_b are feasible. Finally, if $P_a(x_a, x_b) = 1 < \tilde{P}_a(x_a, x_b)$, then party a 's utility $U_a(x_a, x_b)$ is strictly less than $\tilde{U}_a(x_a, x_b)$, because, by the above, x_a is strictly closer to x_a^* than x_b . Further $\tilde{U}_a(x_a, x_b)$ is strictly less than $\max_{x'_a} \tilde{U}_a(x'_a, x_b) = \tilde{U}_a(\hat{x}_a, x_b) = U_a(\hat{x}_a, x_b)$. Therefore, x_a delivers a lower utility than $\hat{x}_a$.

Finally, consider the remaining possibility that feasibility is binding at (x_a, x_b) , x_a is closer to x_a^* than is x_b , and both parties obtain positive expected vote shares. We distinguish between three subcases, depending on which of the three policies $\{x_a, \hat{x}_a, x_b\}$ is between the other two.

In the first subcase, x_a is between $\hat{x}_a$ and x_b . One can easily verify that $\theta_a(x_a, x_b)$ and $\theta_b(x_a, x_b)$ are between $\theta_a(\hat{x}_a, x_b)$ and $\theta_b(\hat{x}_a, x_b)$. Because $\{\theta_a(\hat{x}_a, x_b), \theta_b(\hat{x}_a, x_b)\}$ is feasible, $\{\theta_a(x_a, x_b), \theta_b(x_a, x_b)\}$ is feasible (if 0 is between $\theta_a(x_a, x_b)$ and $\theta_b(x_a, x_b)$) or one party obtains all the votes (otherwise). We have already ruled out both of possibilities above.

In the second subcase, x_b is between $\hat{x}_a$ and x_a . Because $\hat{x}_a$ delivers higher utility than $x_a = x_b$, $\hat{x}_a$ is closer to x_a^* than is x_b . Thus, this subcase implies that x_a is further away from x_a^* than is x_b ,

which we have ruled out above.

In the third and final subcase, $\hat{x}_a$ is between x_a and x_b . By $|\hat{x}_a| \geq |x_b|$, this implies that $|x_a| \geq |\hat{x}_a| \geq |x_b|$. Thus, if feasibility is binding at (x_a, x_b) , then, by Corollary 3, $P_a(x_a, x_b) < \tilde{P}_a(x_a, x_b)$, so $U_a(x_a, x_b)$ is strictly less than $\tilde{U}_a(x_a, x_b)$, by x_a being closer to x_a^* than is x_b . Further, $\tilde{U}_a(x_a, x_b)$ is strictly less than $\tilde{U}_a(\hat{x}_a, x_b) = U_a(\hat{x}_a, x_b)$. Therefore, x_a delivers a lower utility than $\hat{x}_a$.

By symmetry, analogous statements hold when exchanging a and b . $\square$

B.1.6 Electoral Competition under Symmetric Setup

Proposition 5. *Let $x_a^* < 0 < x_b^* = -x_a^*$. The following constitutes an equilibrium.*

	Party platforms	Revealed voter ideology
$\kappa \geq 1$:	Downsian equilibrium $(x_a, x_b) = (0, 0)$	$\rho = \delta(0)$
$\kappa \in (\underline{\kappa}, 1)$:	Polarizing equilibrium $(x_a, x_b) = (1 - \kappa)(x_a^*, x_b^*)$	$\rho = \frac{1}{2}\delta\left(\frac{x_a}{\kappa}\right) + \frac{1}{2}\delta\left(\frac{x_b}{\kappa}\right)$

This equilibrium is unique for $\underline{\kappa} \leq \kappa < 1/2$ and for $\kappa > 1$. For $1/2 < \kappa < 1$, the equilibrium is unique if $|x_a^| \leq \mathbb{E}[|\theta|]$.³⁷*

Proof. First, we verify that the candidate equilibria are in fact equilibria. Second, we show uniqueness.

Solving for a solution to the equations $x_a = x_a^* + \frac{\kappa}{1-\kappa}x_b$ and $x_b = x_b^* + \frac{\kappa}{1-\kappa}x_a$, we obtain

$$(x_a, x_b) = (1 - \kappa)(x_a^*, x_b^*) \quad (32)$$

if $\kappa \in (0, 1)$ and $\kappa \neq 1/2$. By $\kappa > \underline{\kappa}$, the induced distribution over posterior means is feasible, so feasibility is not binding at (x_a, x_b) . By $|x_a| = |x_b|$ and the sufficiency part of Lemma 12, this constitutes mutual best responses and hence an equilibrium.

Uniqueness By the necessity part of Lemma 12 and the solution to (32) being unique, there cannot be another equilibrium (x_a, x_b) where feasibility is not binding if $\kappa \neq 1/2$. Also, we have ruled out (by our refinement) equilibria where one party obtains zero votes. Thus, *it remains to show that there are no equilibria (x_a, x_b) where feasibility is binding and both parties obtain positive expected vote shares.* Denote the equilibrium positions by $(\tilde{x}_a, \tilde{x}_b) = (1 - \kappa)(x_a^*, x_b^*)$.

First, we show that it cannot be the case in equilibrium that both parties propose policies further away from the origin than $\tilde{x}_a$ and $\tilde{x}_b$, respectively, that is, $x_a < \tilde{x}_a < \tilde{x}_b < x_b$. Suppose without loss that x_b is weakly closer to 0 than is x_a . Then, by feasibility being binding and Corollary 3, $P_a(x_a, x_b) \leq \tilde{P}_a(x_a, x_b)$. Because in any equilibrium with positive expected vote shares, x_a must be weakly closer to x_a^* than is x_b , this implies that $U_a(x_a, x_b) \leq \tilde{U}_a(x_a, x_b)$. By $x_a \leq -x_b < \tilde{x}_a$, $-x_b$ is closer to $\hat{x}_a = x_a^* + \frac{\kappa}{1-\kappa}x_b > \tilde{x}_a$ than is x_a . So, $\tilde{U}_a(x_a, x_b) < \tilde{U}_a(-x_b, x_b)$. Finally, $\tilde{U}_a(-x_b, x_b) =$

³⁷We conjecture that this condition can be dropped. Further, we conjecture that the additional equilibria under the knife-edge case $\kappa = 1/2$ cannot be approximated via an arbitrarily small office benefit.

$U_a(-x_b, x_b)$ because $P_a(-x_b, x_b) = \tilde{P}_a(-x_b, x_b) = 1/2$. Together, $U_a(x_a, x_b) < U_a(-x_b, x_b)$ and (x_a, x_b) is not an equilibrium.

Second, we show that there cannot be an equilibrium where $x_a > \tilde{x}_a$ or $x_b < \tilde{x}_b$. To show this, we distinguish between $\kappa > 1/2$ and $\kappa < 1/2$, which affects whether best response of a party, according to Lemma 12, is responsive to other party's position less or more than 1-to-1 by $\frac{\kappa}{1-\kappa} > 1$ and $\frac{\kappa}{1-\kappa} < 1$, respectively.

Case 1: $1/2 < \kappa < 1$.

For the sake of contradiction, suppose $x_b < \tilde{x}_b$. The proof for $x_a > \tilde{x}_a$ is analogous.

If $(1 - \kappa)x_a^* \leq x_b \leq \tilde{x}_b$, then $\hat{x}_a = x_a^* + \frac{\kappa}{1-\kappa}x_b$ is the best response of party a by the following. First, it can be shown that under these conditions, $\tilde{P}_a(\hat{x}_a, x_b) > 0$, so $\theta_a(\hat{x}_a, x_b) \leq 0 \leq \theta_b(\hat{x}_a, x_b)$. Further, the distance $|x_b - \hat{x}_a|$ increases as x_b decreases by $\frac{\kappa}{1-\kappa} > 1$. The maximal distance between x_b and $\hat{x}_a$ is thereby reached at $x_b = (1 - \kappa)x_a^*$, in which case $\hat{x}_a = (1 + \kappa)x_a^*$ and $|x_b - \hat{x}_a| = 2\kappa x_a^*$. By

$$2x_a^* \leq \Delta \Rightarrow 2\kappa x_a^* \leq \kappa\Delta,$$

and Lemma 13, feasibility is not binding at $(\hat{x}_a, x_b)$. Further, by $|\hat{x}_a| \geq |x_b|$ and Lemma 12, $\hat{x}_a$ is the best response of party a to x_b . On the other hand, if $x_b < (1 - \kappa)x_a^*$, then the best response x_a must be closer to x_a^* than x_b is, so $|x_b - x_a| < 2\kappa x_a^*$ as well. By the above, feasibility is not binding at (x_a, x_b) again. Thus, feasibility is not binding in the equilibrium, which we have ruled out above.

Case 2: $0 < \kappa < 1/2$.

Suppose party a chose a position x_a in the direction of the origin from $\tilde{x}_a$, $\tilde{x}_a < x_a$. Define $\Delta := x_a - \tilde{x}_a$ and $\hat{x}_b := x_b^* + \frac{\kappa}{1-\kappa}x_a$. Then, by $\frac{\kappa}{1-\kappa} < 1$, $\hat{x}_b - \tilde{x}_b < \Delta$. By $x_a < x_b^*$ (Lemma 8), we have $x_a \leq \hat{x}_b$, so together, $|\hat{x}_b - x_a| < |\tilde{x}_b - \tilde{x}_a|$. By Jewitt's Lemma, feasibility is not binding at $(x_a, \hat{x}_b)$ or party a obtains all the votes. In the former case, $\hat{x}_b$ is the best response of party b by $|\hat{x}_b| > |x_a|$ and the sufficiency part of Lemma 12, so feasibility would not be binding at the equilibrium. In the latter case, the best response of party b would give party b no votes. We have ruled out both possibilities above. Analogously, we can rule out $x_b < \tilde{x}_b$ in equilibrium.

Lemma 10 concludes the proof. $\square$

B.1.7 Electoral Competition under Asymmetric Setup

Proposition 6. Suppose $x_a^* < 0 < x_b^*$ and $|x_a^*| \leq |x_b^*|$. The positions

$$\begin{aligned} x_a(\kappa) &= \frac{1 - \kappa}{1 - 2\kappa} \left((1 - \kappa)x_a^* + \kappa x_b^* \right) \\ x_b(\kappa) &= \frac{1 - \kappa}{1 - 2\kappa} \left((1 - \kappa)x_b^* + \kappa x_a^* \right) \end{aligned}$$

constitute an equilibrium for

$$\kappa \in \left(0, \frac{x_b^* - x_a^*}{3x_b^* - x_a^*} \right) \cup \left(\frac{x_b^* - x_a^*}{x_b^* - 3x_a^*}, 1 \right) \quad \text{and} \quad \max\{|x_a(\kappa)|, |x_b(\kappa)|\} < \kappa \mathbb{E}[|\theta|]. \quad (33)$$

This is the unique equilibrium for κ satisfying (33) if additionally $|x_b^* - x_a^*| < \kappa \mathbb{E}[|\theta|]$.

If $\kappa \geq 1$, the unique equilibrium is $(x_a, x_b) = (0, 0)$.

Proof. Define

$$\mathcal{K} = \left\{ \kappa \in \left(0, \frac{x_b^* - x_a^*}{3x_b^* - x_a^*} \right) \cup \left(\frac{x_b^* - x_a^*}{x_b^* - 3x_a^*}, 1 \right) \mid \max\{|x_a(\kappa)|, |x_b(\kappa)|\} < \kappa \mathbb{E}[\theta | \theta > 0] \right\}.$$

First, we verify that $(x_a(\kappa), x_b(\kappa))$ constitutes an equilibrium for $\kappa \in \mathcal{K}$. Then, we show uniqueness.

Equilibrium verification If $\kappa \geq 1$, then the unique equilibrium is $(x_a, x_b) = (0, 0)$ by Lemma 10.

Suppose $\kappa < 1$. Using the mutual best responses

$$\begin{aligned} x_a &= x_a^* + \frac{\kappa}{1 - \kappa} x_a \\ x_b &= x_b^* + \frac{\kappa}{1 - \kappa} x_b, \end{aligned}$$

one can solve for the unique solution

$$\begin{aligned} \tilde{x}_a &= \frac{1 - \kappa}{1 - 2\kappa} \left((1 - \kappa)x_a^* + \kappa x_b^* \right), \\ \tilde{x}_b &= \frac{1 - \kappa}{1 - 2\kappa} \left((1 - \kappa)x_b^* + \kappa x_a^* \right). \end{aligned}$$

if $\kappa \neq 1/2$.

Part 1: We verify first that feasibility is not binding at $(\tilde{x}_a, \tilde{x}_b)$.

Under $\kappa > 1/2$, $\tilde{x}_a$ is more extreme than $\tilde{x}_b$ by $|x_b^*| > |x_a^*|$. The posterior means $\theta_a(\tilde{x}_a, \tilde{x}_b)$ and $\theta_b(\tilde{x}_a, \tilde{x}_b)$ are on opposite sides of the origin if

$$0 < \theta_b(\tilde{x}_a, \tilde{x}_b) = \frac{\tilde{x}_a + \tilde{x}_b}{2} + \frac{\tilde{x}_b - \tilde{x}_a}{2\kappa} \Rightarrow \kappa > \frac{x_b^* - x_a^*}{x_b^* - 3x_a^*} > \frac{1}{2}.$$

Under $\kappa < 1/2$, $\tilde{x}_b$ is more extreme than $\tilde{x}_a$ by $|x_b^*| > |x_a^*|$. The posterior means $\theta_a(\tilde{x}_a, \tilde{x}_b)$ and $\theta_b(\tilde{x}_a, \tilde{x}_b)$ are on opposite sides of the origin if

$$0 > \theta_a(\tilde{x}_a, \tilde{x}_b) = \frac{\tilde{x}_a + \tilde{x}_b}{2} - \frac{\tilde{x}_b - \tilde{x}_a}{2\kappa} \Rightarrow \kappa < \frac{x_b^* - x_a^*}{3x_b^* - x_a^*} < \frac{1}{2}.$$

It remains to show that the prior is wide enough, so $\{\theta_a, \theta_b\}$ are feasible.

Define

$$\Delta := \mathbb{E}[\theta | \theta > 0] - \mathbb{E}[\theta | \theta < 0].$$

Lemma 13. If $|x_b - x_a| \leq \kappa \Delta$ and $\theta_a(x_a, x_b)$ and $\theta_b(x_a, x_b)$ are on different sides of the origin, then feasibility is not binding at (x_a, x_b) .

Proof. By Jewitt's Lemma, if θ_a and θ_b are on different sides of the origin and $|\theta_b - \theta_a| < \Delta$, then

$\{\theta_a, \theta_b\}$ is feasible.

$$|\theta_b - \theta_a| = \frac{|x_b - x_a|}{\kappa} < \Delta \Leftrightarrow |x_b - x_a| < \Delta.$$

□

We have

$$\tilde{x}_b - \tilde{x}_a = (1 - \kappa)(x_b^* - x_a^*).$$

Thus, as long as

$$(1 - \kappa)(x_b^* - x_a^*) < \kappa\Delta \Leftrightarrow \kappa > \kappa^* = \frac{\frac{x_b^* - x_a^*}{\Delta}}{1 + \frac{x_b^* - x_a^*}{\Delta}},$$

feasibility is not binding at $(\tilde{x}_a, \tilde{x}_b)$ if $\theta_a(\tilde{x}_a, \tilde{x}_b) \leq 0 \leq \theta_b(\tilde{x}_a, \tilde{x}_b)$.

Part 2: It remains to show that $\tilde{x}_a$ and $\tilde{x}_b$ are mutual best responses when $\kappa \in (\kappa^*, \frac{x_b^* - x_a^*}{3x_b^* - x_a^*}) \cup (\frac{x_b^* - x_a^*}{x_b^* - 3x_a^*}, 1)$.

For mutual best responses, by Lemma 12, the party with the more extreme policy platform (party a if $\kappa > 1/2$ and party b if $\kappa < 1/2$) is playing a best response. To show that the party with the less extreme policy platform, say a , is playing a best response, it is sufficient that feasibility is not binding at $(-x_b, x_b)$, by the following lemma.

Lemma 14. *Let $\hat{x}_a = x_a^* + \frac{\kappa}{1-\kappa}x_b$. If $\hat{x}_a < x_b$ and feasibility is not binding at $(\hat{x}_a, x_b)$ nor $(-x_b, x_b)$, then $\hat{x}_a$ is a best response to x_b .*

Proof. We know that $\hat{x}_a$ is a best response among the x_a such that feasibility is not binding at (x_a, x_b) . We rule out step-by-step other x_a .

Any $x_a \geq x_b$ is suboptimal because it is further from x_a^* than x_b and $U_a(\hat{x}_a, x_b) > u(x_b, x_a^*)$ by (31).

Any x_a with $\hat{x}_a < x_a < x_b$ is suboptimal by the following. The posterior means $\theta_a(x_a, x_b)$ and $\theta_b(x_a, x_b)$ are between $\theta_a(\hat{x}_a, x_b)$ and $\theta_b(\hat{x}_a, x_b)$. So, either feasibility is not binding, which we have ruled out above, or party a obtains an expected vote share of 0 or 1. An expected vote share of 0 delivers utility $u(x_b, x_a^*)$, which is smaller than $U_a(\hat{x}_a, x_b)$ and therefore suboptimal. If the expected vote share is 1, then $\tilde{P}_a(x_a, x_b) \geq 1$ (by the posterior means (θ_a, θ_b) being feasible if on different sides of 0), so $U_a(x_a, x_b) \leq \tilde{U}_a(x_a, x_b) < \tilde{U}_a(\hat{x}_a, x_b) = U_a(\hat{x}_a, x_b)$, so x_a is suboptimal.

Any x_a with $x_a < \hat{x}_a$ and $|x_a| < |x_b|$ has non-binding feasibility at (x_a, x_b) , and is ruled out above, by the following. If $x_b < 0$, then $x_a < \hat{x}_a < x_b < 0$ implies that $|x_a| > |x_b|$, a contradiction. If, otherwise, $x_b > 0$, then, because feasibility is not binding at $(-x_b, x_b)$ and $(\hat{x}_a, x_b)$ has non-binding feasibility, feasibility must be non-binding at (x_a, x_b) .

The remaining x_a satisfy $x_a < \hat{x}_a$, feasibility is binding at (x_a, x_b) , and $|x_a| > |x_b|$. Then, by Corollary 3, either $P_a(x_a, x_b) = 0$, which we have ruled out above, or $P_a(x_a, x_b) < \tilde{P}_a(x_a, x_b)$. If x_a is further away from x_a^* than x_b is, then x_a is clearly suboptimal. So, x_a is closer to x_a^* than is x_b , so $U_a(x_a, x_b) < \tilde{U}_a(x_a, x_b) < \tilde{U}_a(\hat{x}_a, x_b) = U_a(\hat{x}_a, x_b)$, so x_a is suboptimal. □

This implies the following lemma.

Lemma 15. Let $\hat{x}_a = x_a^* + \frac{\kappa}{1-\kappa}x_b$. If (x_b, κ) and $(\hat{x}_a, \kappa)$ lie in the cone

$$\mathcal{C} = \left\{ (x, \kappa) \mid |x| \leq \kappa \mathbb{E}[\theta | \theta > 0] \right\},$$

and $\theta_a(\hat{x}_a, x_b) \leq 0 \leq \theta_b(\hat{x}_a, x_b)$, then $\hat{x}_a$ is the best response to x_b . An analogous statement holds when replacing a and b .

Proof. Feasibility is not binding at $(-x_b, x_b)$ because $(-x_b, \kappa)$ lies in the cone $\mathcal{C}$ if (x_b, κ) does and $\theta_a(-x_b, x_b) \leq 0 \leq \theta_b(-x_b, x_b)$. If $(\hat{x}_a, \kappa)$ also lies in the cone, then $|x_b - x_a| \leq \kappa\Delta$. If $\theta_a(\hat{x}_a, x_b) \leq 0 \leq \theta_b(\hat{x}_a, x_b)$, then feasibility is not binding at $(\hat{x}_a, x_b)$ by Lemma 13. Then, by Lemma 14, $\hat{x}_a$ is a best response to x_b . $\square$

As long as

$$\max\{|x_a|, |x_b|\} \leq \kappa \mathbb{E}[\theta | \theta > 0],$$

both (x_a, κ) and (x_b, κ) lie in the cone $\mathcal{C}$. If, additionally,

$$\kappa \in \left(0, \frac{x_b^* - x_a^*}{3x_b^* - x_a^*}\right) \cup \left(\frac{x_b^* - x_a^*}{x_b^* - 3x_a^*}, 1\right),$$

then $\theta_a(\hat{x}_a, x_b) \leq 0 \leq \theta_b(\hat{x}_a, x_b)$, so by Lemma 15, x_a and x_b are mutual best responses.

Uniqueness Let (x_a, x_b) be an equilibrium. By Lemma 8 and 9, we have $2x_a^* \leq x_a \leq x_b \leq 2x_b^*$. By assumption, we have $2|x_b^* - x_a^*| \leq \kappa\Delta$, so

$$|x_b - x_a| \leq 2|x_b^* - x_a^*| \leq \kappa\Delta,$$

so by Lemma 13, either feasibility is not binding at (x_a, x_b) or one party obtains all the votes. We have ruled out both cases. so by Lemma 13, either feasibility is not binding at (x_a, x_b) or voters learn nothing. If feasibility is not binding, we have shown that the equilibrium positions must be $(x_a(\kappa), x_b(\kappa))$. If voters learn nothing, then either one party obtains all the votes or $x_a = x_b$. We have ruled out the former by our equilibrium definition. In equilibrium, it cannot be that $x_a = x_b$ by the following. Suppose, that $x_a = x_b \geq 0$ (the other case is analogous). Then, x_a would obtain a positive vote share by choosing $x'_a = x_a - \varepsilon$ for ε small enough. This would increase the utility $U_a(x_a, x_b)$ because the implemented policy is improved with positive probability (but never worse). $\square$

B.2 Multidimensional Policy Space

We prove Theorems 5 and 6 in parallel through multiple steps, building on our results for a one-dimensional policy space. In section B.2.1, we show how the voter learning problem and the electoral competition game reduce, in specific senses, to one dimension. In section B.2.2, we use this reduction and the one-dimensional best response lemma 12, to prove a lemma on party best responses in a multidimensional policy space. In section B.2.3, we verify that our equilibrium candidate is an equilibrium using the sufficiency part of the best response lemma. In section B.2.4,

we prove uniqueness using the sufficiency part of the best response lemma. Finally, in section B.2.5, we consider the case $\kappa \geq 1$.

B.2.1 Reduction to One Dimension

As stated in the main text, the voter's maximization problem can be written, up to a constant, as a function of the distribution $\rho \in \Delta(\mathbb{R}^n)$ over posterior means θ .

$$\begin{aligned} \max_{\rho \in \Delta(\mathbb{R}^n)} \mathbb{E}_{\theta \sim \rho} \left[\max \left\{ \left\langle x_b - x_a, \theta - \frac{x_a + x_b}{2} \right\rangle, -\left\langle x_b - x_a, \theta - \frac{x_a + x_b}{2} \right\rangle \right\} - \kappa \langle \theta, \theta \rangle \right] \\ \text{s.t. } \rho \leq_{\text{MPS}} \mu. \end{aligned} \quad (\text{Pn})$$

We define an equivalent one-dimensional voter learning problem using the following scalar projection on the line through the origin with direction $x_b - x_a$. We extend the scalar projection on $x_b - x_a$, $\text{proj}_{x_b - x_a}$, to probability measures on $\mathbb{R}^n$ via the pushforward. Define

$$\begin{aligned} \hat{x}_a &= \text{proj}_{x_b - x_a}(x_a) \\ \hat{x}_b &= \text{proj}_{x_b - x_a}(x_b) \\ \hat{\mu} &= \text{proj}_{x_b - x_a}(\mu). \end{aligned} \quad (34)$$

Now we can define (P1) as the following one-dimensional learning problem.

$$\begin{aligned} \max_{\hat{\rho} \in \Delta(\mathbb{R})} \mathbb{E}_{\hat{\theta} \sim \hat{\rho}} \left[\max \left\{ (\hat{x}_b - \hat{x}_a) \left(\hat{\theta} - \frac{\hat{x}_a + \hat{x}_b}{2} \right), -(\hat{x}_b - \hat{x}_a) \left(\hat{\theta} - \frac{\hat{x}_a + \hat{x}_b}{2} \right) \right\} - \kappa \hat{\theta}^2 \right] \\ \text{s.t. } \hat{\rho} \leq_{\text{MPS}} \hat{\mu} \end{aligned} \quad (\text{P1})$$

The following lemma gives a sense in which the two voter learning problems are equivalent.

Lemma 16. *A distribution $\rho \in \Delta(\mathbb{R}^n)$ solves (Pn) if and only if $\text{proj}(\rho)$ solves (P1) and ρ is supported on the line through the origin with the direction $x_b - x_a$.*

Proof. Let $\rho \in \Delta(\mathbb{R}^n)$ be supported on the line through the origin with direction $x_b - x_a$. We show that then the objectives of (Pn) and (P1) are equivalent. An equivalent way of representing $\theta \sim \rho$ is as $\hat{\theta} \frac{x_b - x_a}{\|x_b - x_a\|}$ with $\hat{\theta} \sim \hat{\rho} = \text{proj}(\rho)$. So, using the definitions (34), the objective of (Pn) can be written as

$$\begin{aligned} \mathbb{E}_{\hat{\theta} \sim \hat{\rho}} \left[\max \left\{ \left\langle x_b - x_a, \hat{\theta} \frac{x_b - x_a}{\|x_b - x_a\|} - \frac{x_a + x_b}{2} \right\rangle, -\left\langle x_b - x_a, \hat{\theta} \frac{x_b - x_a}{\|x_b - x_a\|} - \frac{x_a + x_b}{2} \right\rangle \right\} - \kappa \hat{\theta}^2 \right] \\ = \mathbb{E}_{\hat{\theta} \sim \hat{\rho}} \left[\max \left\{ (\hat{x}_b - \hat{x}_a) \left(\hat{\theta} - \frac{\hat{x}_a + \hat{x}_b}{2} \right), -(\hat{x}_b - \hat{x}_a) \left(\hat{\theta} - \frac{\hat{x}_a + \hat{x}_b}{2} \right) \right\} - \kappa \hat{\theta}^2 \right], \end{aligned}$$

which is the objective of (P1).

By the following statement, the constraints $\rho \leq_{\text{MPC}} \mu$ and $\hat{\rho} \leq_{\text{MPC}} \text{proj}(\mu)$ are equivalent: *Let $\rho \in \Delta(\mathbb{R}^n)$ be supported on the line through the origin with direction $x_b - x_a$. Then $\rho \leq_{\text{MPS}} \mu$ if and only if $\text{proj}(\rho) \leq_{\text{MPS}} \text{proj}(\mu)$.*

The only if-direction follows directly from the fact that the mean-preserving contraction relation is preserved by scalar projections.

The if-direction follows from the fact that ρ is already supported on a line that is preserved by the orthogonal projection associated with the scalar projection proj . Thus, $\text{proj}(\rho) \leq_{\text{MPS}} \text{proj}(\mu)$ directly implies that ρ is a mean-preserving contraction of the orthogonal projection of μ on the line through the origin with direction $x_b - x_a$. The orthogonal projection of μ is a mean-preserving contraction of μ because μ is spherical. By transitivity of mean-preserving contraction, $\text{proj}(\rho) \leq_{\text{MPS}} \mu$.

Collecting results, ρ solving (Pn) and $\text{proj}(\rho)$ solving (P1) are equivalent for distributions ρ supported on the line through the origin with direction $x_b - x_a$. This establishes the if direction. For the only if direction, Theorem 1 implies that any solution ρ to (Pn) is supported on the line through the origin with direction $x_b - x_a$. $\square$

This reduction of the voter learning problem to one dimension allows us to apply results on voter learning under a one-dimensional policy space. We generalize the definitions of feasibility and binding feasibility (Definition 1 and 2) verbatim to the multidimensional policy space using the definition of $\theta_a(x_a, x_b)$ and $\theta_b(x_a, x_b)$ from (11). Proposition 3 follows from Corollary 4. By Corollary 3, the vote share for party a is

$$P_a(x_a, x_b) = \frac{1}{2} + \frac{\kappa \hat{x}_a + \hat{x}_b}{2 \hat{x}_b - \hat{x}_a} = \frac{1}{2} + \frac{\kappa \|x_b\|^2 - \|x_a\|^2}{2 \|x_b - x_a\|^2}$$

if feasibility is not binding at (x_a, x_b) .

Now, we can show that also for parties there is an equivalence to a one-dimensional game, up to constants, as long as their positions are on a certain line. The following lemma shows that the party objectives are the same as if we consider the one-dimensional game where party positions and ideal policies are projected on the line through the party positions x_a and x_b . Thus, if we restrict party positions to lie on some line through the policy space, we can solve the game like the game with a one-dimensional policy space. The result follows essentially from the quadratic preferences of parties and the one-dimensional reduction of the voter learning problem.

Recall that each party j chooses its position $x_j \in \mathbb{R}^n$ to maximize its objective

$$U_j(x_a, x_b) = P_a(x_a, x_b)u(x_a, x_j^*) + (1 - P_a(x_a, x_b))u(x_b, x_j^*),$$

where $P_a(x_a, x_b)$ is the mass of voter revealed ideal points that are closer to x_a than to x_b , as determined by optimal voter learning.

To introduce the equivalent one-dimensional model, we need the following definitions. Using the same scalar projection proj on the line through the party positions, defined in (12), we define $\hat{x}_a, \hat{x}_b$ as in (34) and, additionally,

$$\begin{aligned} \hat{x}_a^* &= \text{proj}_{x_b - x_a}(x_a^*) \\ \hat{x}_b^* &= \text{proj}_{x_b - x_a}(x_b^*). \end{aligned}$$

Party a 's objective in a one-dimensional policy space for $\hat{x}_a, \hat{x}_b \in \mathbb{R}$ is

$$\begin{aligned}\hat{U}_a(\hat{x}_a, \hat{x}_b) &:= \hat{P}_a(\hat{x}_a, \hat{x}_b) \hat{u}(\hat{x}_a, \hat{x}_a^*) + (1 - \hat{P}_a(\hat{x}_a, \hat{x}_b)) \hat{u}(\hat{x}_b, \hat{x}_a^*) \\ &= \hat{P}_a(\hat{x}_a, \hat{x}_b) 2(\hat{x}_a - \hat{x}_b) \left(\hat{x}_a^* - \frac{\hat{x}_a + \hat{x}_b}{2} \right) + u(\hat{x}_b, \hat{x}_a^*)\end{aligned}$$

where $\hat{u}(x, y) = -(x - y)^2$ and here $\hat{P}_a(\hat{x}_a, \hat{x}_b)$ is the mass of voter revealed ideal points that are closer to $\hat{x}_a$ than to $\hat{x}_b$, as determined by optimal voter learning under a one-dimensional policy space. Finally, let D^2 be the squared distance of x_a^* from the line through x_a and x_b .

The following lemma implies the party objective is the same as in the one-dimensional model up to a constant D^2 that only depends on x_a and x_b through the line on which they are.

Lemma 17. *The party utility $U_a(x_a, x_b)$ satisfies*

$$U_a(x_a, x_b) = \hat{U}_a(\hat{x}_a, \hat{x}_b) + D^2 = \hat{P}_a(\hat{x}_a, \hat{x}_b) 2(\hat{x}_a - \hat{x}_b) \left(\hat{x}_a^* - \frac{\hat{x}_a + \hat{x}_b}{2} \right) + u(x_b, x_a^*).$$

An analogous identity holds after swapping a and b .

Proof. By Lemma 16, $P_a(x_a, x_b) = \hat{P}_a(\hat{x}_a, \hat{x}_b)$. The party objective is

$$\begin{aligned}U_a(x_a, x_b) &= P_a(x_a, x_b) (u(x_a, x_a^*) - u(x_b, x_b^*)) + u(x_b, x_a^*) \\ &= \hat{P}_a(\hat{x}_a, \hat{x}_b) 2 \left\langle x_a - x_b, x_a^* - \frac{x_a + x_b}{2} \right\rangle + u(x_b, x_a^*) \\ &= \hat{P}_a(\hat{x}_a, \hat{x}_b) 2(\hat{x}_a - \hat{x}_b) \left(\hat{x}_a^* - \frac{\hat{x}_a + \hat{x}_b}{2} \right) + \hat{u}(\hat{x}_b, \hat{x}_a^*) + D^2 \\ &= \hat{U}_a(\hat{x}_a, \hat{x}_b) + D^2\end{aligned}$$

where we have used Pythagoras theorem in the middle step. □

B.2.2 Multidimensional Best Response Lemma

The following lemma on best responses exploits the reduction to one dimension proven by the above lemma. Following definition (13), let $P_{(1-\kappa)x_a^*, x_b}$ denote the orthogonal projection of $\mathbb{R}^n$ on the line through $(1-\kappa)x_a^*$ and x_b . Recall that we generalize the definitions of feasibility and binding feasibility (Definition 1 and 2) verbatim to the multidimensional policy space using (11). Further, we say feasibility is strictly non-binding at (x_a, x_b) if there exists open neighborhoods $\mathcal{N}_a, \mathcal{N}_b$ of x_a and x_b , respectively, such that feasibility is not binding for all (x'_a, x'_b) with $x'_a \in \mathcal{N}_a$ and $x'_b \in \mathcal{N}_b$.

Lemma 18 (Multidimensional Best Response Lemma). *Let $\kappa < 1$, $x_b \in \mathbb{R}^n$. Define*

$$\tilde{x}_a := P_{(1-\kappa)x_a^*, x_b} \left(x_a^* + \frac{\kappa}{1-\kappa} x_b \right).$$

Necessity: *If x_a is the best response of party a to x_b and feasibility is strictly non-binding at (x_a, x_b) , then $x_a = \tilde{x}_a$.*

Sufficiency: *If feasibility is not binding at $(\tilde{x}_a, x_b)$ or $(-x_b, x_b)$, then $\tilde{x}_a$ is the best response to x_b .*

Proof. Necessity: Suppose x_a is a best response to x_b and feasibility is strictly non-binding at (x_a, x_b) .

By Lemma 17,

$$U_a(x_a, x_b) = \hat{P}_a(\hat{x}_a, \hat{x}_b)2(\hat{x}_a - \hat{x}_b)\left(\hat{x}_a^* - \frac{\hat{x}_a + \hat{x}_b}{2}\right) + u(x_b, x_a^*),$$

where $\hat{x}_a, \hat{x}_b, \hat{x}_a^*$ are again the scalar projections of x_a, x_b, x_a^* on the line through x_a and x_b . If feasibility is not binding at (x_a, x_b) , then the vote share $\hat{P}_a(x_a, x_b)$ equals the pseudo vote-share

$$\tilde{P}_a(x_a, x_b) = \frac{1}{2} + \frac{\kappa}{2} \frac{\hat{x}_a + \hat{x}_b}{\hat{x}_b - \hat{x}_a},$$

so the utility $U_a(x_a, x_b)$ equals the pseudo utility

$$\begin{aligned} \tilde{U}_a(x_a, x_b) &:= \tilde{P}_a(x_a, x_b)2(\hat{x}_a - \hat{x}_b)\left(\hat{x}_a^* - \frac{\hat{x}_a + \hat{x}_b}{2}\right) + u(x_b, x_a^*) \\ &= \left(\frac{1}{2} + \frac{\kappa}{2} \frac{\hat{x}_a + \hat{x}_b}{\hat{x}_b - \hat{x}_a}\right)2(\hat{x}_a - \hat{x}_b)\left(\hat{x}_a^* - \frac{\hat{x}_a + \hat{x}_b}{2}\right) + u(x_b, x_a^*) \\ &= ((1 - \kappa)\hat{x}_a - (1 + \kappa)\hat{x}_b)\left(\hat{x}_a^* - \frac{\hat{x}_a + \hat{x}_b}{2}\right) + u(x_b, x_a^*) \end{aligned} \quad (35)$$

If feasibility is strictly non-binding at (x_a, x_b) , then $U_a(x_a, x_b)$ is locally given by the pseudo utility. If x_a is a best response, then x_a must be a local maximum of the pseudo utility $\tilde{U}_a(x_a, x_b)$. Then, the projection $\hat{x}_a = \text{proj}(x_a)$ must satisfy the first-order condition of (35). We know from Lemma 12 that the unique $\hat{x}_a$ that satisfies the first-order condition is $\hat{x}_a = \hat{x}_a^* + \frac{\kappa}{1 - \kappa}\hat{x}_b$. Inserting $\hat{x}_a$ and simplifying, we obtain

$$\begin{aligned} U_a(x_a, x_b) &= \left((1 - \kappa)\left(\hat{x}_a^* + \frac{\kappa}{1 - \kappa}\hat{x}_b\right) - (1 + \kappa)\hat{x}_b\right)\left(\hat{x}_a^* - \frac{\hat{x}_a^* + \frac{\kappa}{1 - \kappa}\hat{x}_b + \hat{x}_b}{2}\right) + u(x_b, x_a^*) \\ &= \left((1 - \kappa)\hat{x}_a^* - \hat{x}_b\right)\frac{1}{2}\left(\hat{x}_a^* - \frac{\hat{x}_b}{1 - \kappa}\right) + u(x_b, x_a^*) \\ &= \frac{1}{2(1 - \kappa)}(\hat{x}_b - (1 - \kappa)\hat{x}_a^*)^2 + u(x_b, x_a^*) \\ &= \frac{1}{2(1 - \kappa)}\frac{\langle x_b - x_a, x_b - (1 - \kappa)x_a^* \rangle^2}{\|x_b - x_a\|^2} + u(x_b, x_a^*). \end{aligned} \quad (36)$$

If x_a is a local maximum of the pseudo utility, then the direction $\frac{x_a - x_b}{\|x_a - x_b\|}$ of x_a from x_b , given that the projection $\hat{x}_a$ on this direction is optimally chosen, must be a local optimum of (36). The only local optimum is when $x_b - x_a$ is parallel to $x_b - (1 - \kappa)x_a^*$, that is, when x_a is on the line through $(1 - \kappa)x_a^*$ and x_b .

Together, we have shown that x_a is on the line through $(1 - \kappa)x_a^*$ and x_b and the scalar projection of x_a on this line is $\hat{x}_a^* + \frac{\kappa}{1 - \kappa}\hat{x}_b$. The unique point that satisfies these two conditions is $\tilde{x}_a$.

Sufficiency: We know from (36) if feasibility is not binding at $\tilde{x}_a$, then

$$U_a(\tilde{x}_a, x_b) = \frac{1}{2(1-\kappa)} \|x_b - (1-\kappa)x_a^*\|^2 + u(x_b, x_a^*).$$

For any $d \in \mathbb{R}^n$ with $\|d\| = 1$, define

$$B(d) := \frac{1}{4(1-\kappa)} (\langle d, x_b - (1-\kappa)x_a^* \rangle)^2 + u(x_b, x_a^*),$$

which is bounded from above by $U_a(\tilde{x}_a, x_b)$. The following lemma concludes the proof.

Lemma 19. *For any $x_a \in \mathbb{R}^n$, the utility $U_a(x_a, x_b)$ is bounded from above by*

$$B\left(\frac{x_b - x_a}{\|x_b - x_a\|}\right).$$

Proof. First note that by (36) the pseudo utility

$$\tilde{U}_a(x_a, x_b) = \tilde{P}_a(x_a, x_b) 2(\hat{x}_a - \hat{x}_b) \left(\hat{x}_a^* - \frac{\hat{x}_a + \hat{x}_b}{2} \right) + u(x_b, x_a^*),$$

is bounded from above by $B(\frac{x_b - x_a}{\|x_b - x_a\|})$.

Let $d \in \mathbb{R}^n$ with $\|d\| = 1$. We show that among all x_a with the same $d(x_a) = \frac{x_b - x_a}{\|x_b - x_a\|}$, no x_a obtains a higher utility than $B(\frac{x_b - x_a}{\|x_b - x_a\|})$.

If x_a is at least as far from x_a^* than x_b , then $U_a(x_a, x_b) \leq u(x_b, x_a^*) \leq B(d)$. Suppose for the rest of the proof that x_a is strictly closer to x_a^* than x_b .

If feasibility is not binding at x_a , then the vote share equals the pseudo vote share, so the utility $U_a(x_a, x_b)$ equals the pseudo utility, which we have argued is bounded from above by $B(d)$.

If feasibility is binding at (x_a, x_b) , then by Proposition 4, either the origin does not lie between the maxima θ_a and θ_b of the symmetrized valued function, or voters acquire a threshold signal. If the origin does not lie between θ_a and θ_b , then either a receives all votes or b receives all votes. If a receives all votes, then $P_a(x_a, x_b) = 1 \leq \tilde{P}_a(x_a, x_b)$, so the utility is bounded from above by the pseudo utility. If b receives all votes, then $U_a(x_a, x_b) = u(x_b, x_a^*)$ which is not greater than $B(d)$. Finally, consider the case that voters acquire a threshold signal. Because feasibility not binding at $(-x_b, x_b)$, by Jewitt's Lemma, x_a must be further away from 0 than x_b . Thus, $\hat{P}_a(x_a, x_b) < \tilde{P}_a(x_a, x_b)$, by Corollary 3, case (2). Again, the utility is bounded from above by the pseudo utility. $\square$

This concludes the proof of sufficiency. $\square$

B.2.3 Equilibrium Verification

Using the best response lemma, we verify that the polarizing equilibria of Theorem 5 and Theorem 6 are in fact equilibria.

First, we turn to the equilibrium under the symmetric setup. If $\kappa < 1$, $\|x_a^*\| = \|x_b^*\|$, and $\frac{1-\kappa}{\kappa} \|x_a^*\| \leq \mathbb{E}[\|\theta\|]$, then $(x_a, x_b) = (1-\kappa)(x_a^*, x_b^*)$ is an equilibrium by the following. If $x_b =$

$(1 - \kappa)x_b^*$ and $\|x_a^*\| = \|x_b^*\|$, then

$$P_{(1-\kappa)x_a^*, x_b} \left(x_a^* + \frac{\kappa}{1-\kappa} x_b \right) = P(x_a^* + \kappa x_b^*) = (1 - \kappa)x_a^*.$$

If $\frac{1-\kappa}{\kappa} \|x_a^*\| \leq \mathbb{E}[\|\theta\|]$, then feasibility is not binding at $(1 - \kappa)(-x_b^*, x_b^*)$ or $(1 - \kappa)(x_a^*, x_b^*)$. Together, by Lemma 18, $x_a = (1 - \kappa)x_a^*$ is a best response to $x_b = (1 - \kappa)x_b^*$. By symmetry, x_b is a best response to x_a .

Next, we turn to the equilibrium under the asymmetric setup. One can verify that the positions $(x_a(\kappa), x_b(\kappa))$ satisfy

$$x_j = P_{(1-\kappa)x_j^*, x_{-j}} \left(x_j^* + \frac{\kappa}{1-\kappa} x_{-j} \right)$$

from the best response lemma, using Proposition 6. If $\max\{\|x_a(\kappa)\|, \|x_b(\kappa)\|\} \leq \kappa \mathbb{E}[\|\theta\|]$, then feasibility is not binding at $(1 - \kappa)(-x_b^*, x_b^*)$ or $(1 - \kappa)(x_a^*, x_b^*)$. So, by the sufficiency part of the best response lemma, $(x_a(\kappa), x_b(\kappa))$ is an equilibrium.

B.2.4 Uniqueness

We prove uniqueness in three steps, under

$$\|x_b^* - x_a^*\| \leq \kappa \mathbb{E}[\|\theta\|].$$

First, we establishing bounds on x_a and x_b that must hold in any equilibrium (x_a, x_b) . Second, show that given these bounds, feasibility is not binding (x_a, x_b) if (x_a, x_b) is an equilibrium. Third, we use the necessity part of the multidimensional best response lemma to reduce the problem to one dimension and then apply our uniqueness results under a one-dimensional policy space.

First Step: Bounds

Let (x_a, x_b) be an equilibrium. We prove the following bound on party polarization:

$$\|x_b - x_a\| \leq 2\|x_b^* - x_a^*\|.$$

If (x_a, x_b) is an equilibrium, then (x_a, x_b) is also an equilibrium when both parties are restricted to positions on the line through x_a and x_b . By Lemma 17, we can use bounds on equilibrium platforms under a one-dimensional policy space. Let $x_a \neq x_b$ (otherwise, we are done) and define

$$\begin{aligned} \hat{x}_a &:= \text{proj}_{x_b - x_a}(x_a) \\ \hat{x}_b &:= \text{proj}_{x_b - x_a}(x_b) \\ \hat{x}_a^* &:= \text{proj}_{x_b - x_a}(x_a^*) \\ \hat{x}_b^* &:= \text{proj}_{x_b - x_a}(x_b^*). \end{aligned}$$

In an equilibrium with positive expected vote shares, we have

$$\begin{aligned} ||x_a - x_a^*||^2 &\leq ||x_b - x_a^*||^2 \\ ||x_b - x_b^*||^2 &\leq ||x_a - x_b^*||^2. \end{aligned}$$

Adding up these equations and collecting terms, we obtain $\hat{x}_a \leq \hat{x}_b$.

These are two cases. Either $\hat{x}_a^*$ and $\hat{x}_b^*$ are on different sides of 0 or they are on the same side. By symmetry, for the second case we can restrict attention without loss to $0 < \hat{x}_a^*, 0 < \hat{x}_b^*$.

Case 1: $\hat{x}_a^* \leq 0 \leq \hat{x}_b^*$.

By Lemmas 8 and 9,

$$2\hat{x}_a^* \leq \hat{x}_a \leq \hat{x}_b \leq 2\hat{x}_b^*.$$

This implies

$$||x_b - x_a|| = |\hat{x}_b - \hat{x}_a| \leq 2|\hat{x}_b^* - \hat{x}_a^*| \leq 2||x_b^* - x_a^*||. \quad (37)$$

Case 2: $0 < \hat{x}_a^*, 0 < \hat{x}_b^*$.

By Lemma 11, $\hat{x}_a = \hat{x}_b$ or

$$\hat{x}_a^* \leq \hat{x}_a \leq \hat{x}_b \leq 2\hat{x}_b^* - \hat{x}_a^*.$$

This implies

$$||x_b - x_a|| = |\hat{x}_b - \hat{x}_a| \leq 2|\hat{x}_b^* - \hat{x}_a^*| \leq 2||x_b^* - x_a^*||.$$

Step 2: Feasibility not binding

In both cases above, we have $||x_a - x_b|| \leq 2||x_b^* - x_a^*||$. Recall that $\frac{||x_b - x_a||}{\kappa}$ is the distance between the posterior means θ_a and θ_b .

By assumption, $||x_b^* - x_a^*|| \leq \kappa \mathbb{E}[||\theta||]$, therefore

$$||x_b - x_a|| \leq 2||x_b^* - x_a^*|| \leq 2\kappa \mathbb{E}[||\theta||] \Rightarrow \frac{||x_b - x_a||}{\kappa} \leq 2\mathbb{E}[||\theta||].$$

The distance between the posterior means of symmetric threshold signal is $2\mathbb{E}[||\theta||]$. By Jewitt's lemma, the distance between the posterior means of any threshold signal is at most $2\mathbb{E}[||\theta||]$. Thus, the mean-preserving contraction constraint on voter learning is not binding in equilibrium.

Feasibility can thus be binding only if voters learn nothing. This implies that one party obtains all the votes (which we have ruled out by our equilibrium definition) or $x_a = x_b$. We rule out $x_a = x_b$ by the following.

First, $(x_a, x_b) = (0, 0)$ is not an equilibrium under $\kappa < 1$. We have assumed that $x_a^* \neq x_b^*$, thus at least one party's ideal policy is not 0. Without loss, assume $x_a^* \neq 0$. Then, party a could profitably deviate by moving toward its own ideal policy, which under $\kappa < 1$ delivers a positive vote share.

Second, we rule out $(x_a, x_b) = x \neq 0$ using the following lemma.

Lemma 20. *If $x \neq 0$ is not on a line through 0, x_a^* , and x_b^* , then (x, x) is not an equilibrium.*

Proof. Let $x_a = x_b = x \neq 0$. We show that (x_a, x_b) is not an equilibrium.

If x is not on the line through 0 and x_a^* , then in any ε -neighborhood of x there is a point x' that is closer to both 0 and x_a^* (x' can be obtained, for example, by moving x in the direction $\frac{1}{2}(x_a^* - x) + \frac{1}{2}(0 - x)$). For a small enough $\varepsilon > 0$, if a adopted the position x' , this would lead a to obtain all votes and a higher policy utility conditional on winning. Thus, x' would be a profitable deviation for party a .

Analogously, if x is not on the line through 0 and x_b^* , b would have a profitable deviation. $\square$

If 0, x_a^* , and x_b^* are not jointly on a line, then x cannot be on a line through all of them, so by Lemma 20, (x, x) is not an equilibrium. Finally, consider the case that x_a^* , 0, and x_b^* are on a line and $x \neq 0$ is on said line. We have assumed that $x_a^* \neq x_b^*$. Thus, for Theorem 5, $\|x_a^*\| = \|x_b^*\|$ implies that 0 is between the projections of x_a^* and x_b^* on said line. Therefore, there is not an equilibrium where $x_a = x_b = x$ by a similar argument as in the proof of Lemma 20: at least one party would profit from moving slightly closer to the origin, obtaining all the votes and moving the implemented policy closer to its ideal policy. For Theorem 6, we have assumed that the scalar projections of x_a^* and x_b^* on the line through x_a^* and x_b^* are on different sides of 0. Thus, by the same argument, there is not an equilibrium where $x_a = x_b = x$.

Third Step: Reduction to one dimension

We have established in the first two steps that in any equilibrium (x_a, x_b) , feasibility is not binding. Then, by the necessity part of the multidimensional best response lemma, x_a and x_b are on the line through $(1 - \kappa)x_a^*$ and $(1 - \kappa)x_b^*$, namely $(x_a, x_b) = (1 - \kappa)(x_a^*, x_b^*)$. By Lemma 17 and Propositions 5 and 6, the equilibrium positions are unique.

B.2.5 Equilibrium if $\kappa \geq 1$

Define

$$\tilde{x}_a^* := \text{proj}_{x_b^* - x_a^*}(x_a^*)$$

$$\tilde{x}_b^* := \text{proj}_{x_b^* - x_a^*}(x_b^*).$$

Lemma 21. *If $\kappa \geq 1$ and $\tilde{x}_a^* < 0 < \tilde{x}_b^*$, then the unique equilibrium is $(0, 0)$.*

Proof. First, $(0, 0)$ is an equilibrium. The reason is that under $\kappa \geq 1$, the implemented policy is 0 if one party chooses position 0. This follows from Corollary 3 and the fact that $\theta_a(x_a, x_b)$ and $\theta_b(x_a, x_b)$ are weakly between x_a and x_b for $\kappa \geq 1$. Hence, no party has a profitable deviation.

Second, we show uniqueness. Suppose, $(x_a, x_b) \neq 0$ is an equilibrium. Define

$$\hat{x}_a^* := \text{proj}_{x_b - x_a}(x_a^*)$$

$$\hat{x}_b^* := \text{proj}_{x_b - x_a}(x_b^*)$$

We distinguish between two cases.

First, consider $\hat{x}_a^* \leq 0 \leq \hat{x}_b^*$. By Lemma 17, we can apply Lemma 10 for $\kappa \geq 1$ to obtain that (x_a, x_b) is not an equilibrium.

Second, consider the case that 0 is not between $\hat{x}_a^*$ and $\hat{x}_b^*$. Without loss, suppose, $0 < \hat{x}_a^*$ and $0 < \hat{x}_b^*$. By Lemma 11, either

$$x_a^* \leq x_a < x_b \leq 2x_b^* - x_a^*$$

or $x_a = x_b$. In the former case, party a obtains all the votes, contradicting our equilibrium definition. In the latter case, by Lemma 20, this is not an equilibrium if $x_a = x_b \neq 0$ if the line through x_a and x_b does not go through x_a^* and x_b^* . The line through x_a and x_b does not go through x_a^* and x_b^* because otherwise 0 would be between $\hat{x}_a^*$ and $\hat{x}_b^*$ by $\tilde{x}_a^* < 0 < \tilde{x}_b^*$. $\square$

C Appendix: IO Extension

C.1 Lemma 2

Proof. In a pure-strategy equilibrium, consumers anticipate product locations, which induces a one-dimensional distribution of revealed preferences (Theorem 1).

Suppose without loss that revealed preferences are supported on the line spanned by

$$e_1 = \begin{pmatrix} 1 \\ 0 \\ \dots \\ 0 \end{pmatrix}.$$

Solving for Equilibrium Profits: First, we show an equivalence between multidimensional product locations $x_a, x_b \in \mathbb{R}^n$ and one-dimensional *pseudo-locations* $\hat{x}_a, \hat{x}_b \in \mathbb{R}$ in the sense that the utility difference is the same,

$$\begin{aligned} u(x_b, \theta) - u(x_a, \theta) &= 2\left\langle \theta - \frac{x_a + x_b}{2}, x_b - x_a \right\rangle_A \\ &= 2\langle x_b - x_a, e_1 \rangle_A \theta_1 + (\langle x_a, x_a \rangle_A - \langle x_b, x_b \rangle_A) \\ &= 2\left(\theta_1 - \frac{\langle x_b, x_b \rangle_A - \langle x_a, x_a \rangle_A}{2\langle x_b - x_a, e_1 \rangle_A}\right) \langle x_b - x_a, e_1 \rangle_A \\ &= 2\left(\theta_1 - \frac{\hat{x}_a + \hat{x}_b}{2}\right) (\hat{x}_b - \hat{x}_a) = \hat{u}(\hat{x}_b, \theta_1) - \hat{u}(\hat{x}_a, \theta_1) \end{aligned}$$

where $\hat{x}_a, \hat{x}_b \in \mathbb{R}$ are uniquely determined by

$$\begin{aligned} \hat{x}_b - \hat{x}_a &= \langle x_b - x_a, e_1 \rangle_A \\ \frac{\hat{x}_a + \hat{x}_b}{2} &= \frac{\langle x_b, x_b \rangle_A - \langle x_a, x_a \rangle_A}{\langle x_b - x_a, e_1 \rangle_A}. \end{aligned}$$

Because the utility difference is a sufficient statistic for subsequent behavior of consumers, this implies behavioral equivalence of consumers. This behavioral equivalence allows us to use results from Anderson, Goeree, and Ramer (1997) to determine the equilibrium of the pricing subgame and characterize equilibrium profits as a function of product locations x_a, x_b .

Assume without loss that x_a is to the left of x_b , that is, $\langle x_a, e_1 \rangle_A < \langle x_b, e_1 \rangle_A$. (Under equality, prices are zero, which is not an equilibrium.) Following Anderson, Goeree, and Ramer (1997), we define $\xi \in \mathbb{R}$ as the indifferent type given $\hat{x}_a, \hat{x}_b \in \mathbb{R}$ (and the resulting equilibrium prices), which is defined implicitly via

$$\xi = \frac{\hat{x}_a + \hat{x}_b}{2} + \frac{1 - 2F(\xi)}{f(\xi)}. \quad (38)$$

As Anderson, Goeree, and Ramer (1997) note, for log-concave F , the term $\frac{1-2F(\xi)}{f(\xi)}$ is non-increasing in ξ and there is a unique solution to (38). Moreover, it follows that the solution is strictly increasing in $\frac{\hat{x}_a + \hat{x}_b}{2}$. They show that the equilibrium price is $p_a = 2(\hat{x}_b - \hat{x}_a)F(\xi)/f(\xi)$ and therefore equilibrium

utility (profits) U_a for firm a are

$$U_a = 2c(\hat{x}_b - \hat{x}_a) \frac{F(\xi)}{f(\xi)} F(\xi).$$

Projecting x_a : Suppose $x_a \in \mathbb{R}^n$ was not on the line of revealed consumer preferences. We show that given any $x_b \in \mathbb{R}^n$, product location x_a is dominated by its projection $\tilde{x}_a = \frac{\langle e_1, x_a \rangle_A}{\langle e_1, e_1 \rangle_A} e_1$ on the line of revealed preferences. Recall that the one-dimensional pseudo-locations $\hat{x}_a, \hat{x}_b \in \mathbb{R}$, which we use to characterize equilibrium profits, are a function of the multidimensional real locations x_a, x_b . To show that equilibrium profits for firm a are higher under $\tilde{x}_a$ than under x_a , we show first note that the difference of the corresponding pseudo-locations is the same,

$$\hat{x}_b(\tilde{x}_a, x_b) - \hat{x}_a(\tilde{x}_a, x_b) = \langle x_b - \tilde{x}_a, e_1 \rangle_A = \langle x_b - x_a, e_1 \rangle_A = \hat{x}_b(x_a, x_b) - \hat{x}_a(x_a, x_b),$$

because $\tilde{x}_a$ and x_a have the same inner product with e_1 by construction of $\tilde{x}_a$. This implies that the implied midpoints of pseudo-locations are ordered,

$$\begin{aligned} \frac{\hat{x}_a(\tilde{x}_a, x_b) + \hat{x}_b(\tilde{x}_a, x_b)}{2} &= \frac{\langle x_b, x_b \rangle_A - \langle \tilde{x}_a, \tilde{x}_a \rangle_A}{\langle x_b - \tilde{x}_a, e_1 \rangle_A} = \frac{\langle x_b, x_b \rangle_A - \langle \tilde{x}_a, \tilde{x}_a \rangle_A}{\langle x_b - x_a, e_1 \rangle_A} \\ &> \frac{\langle x_b, x_b \rangle_A - \langle x_a, x_a \rangle_A}{\langle x_b - x_a, e_1 \rangle_A} = \frac{\hat{x}_a(x_a, x_b) + \hat{x}_b(x_a, x_b)}{2}, \end{aligned}$$

because $\langle x_a, x_a \rangle_A > \langle \tilde{x}_a, \tilde{x}_a \rangle_A$ and $\langle x_b - x_a, e_1 \rangle_A > 0$.

We noted above that ξ is strictly increasing in $\frac{\hat{x}_a + \hat{x}_b}{2}$. Thus, the last equation implies that ξ is larger under $\tilde{x}_a$ and x_b than under x_a and x_b .

This implies that profits U_a are higher under $\tilde{x}_a$ than under x_a . The differentiation of pseudo-locations $\hat{x}_b - \hat{x}_a$ is the same under $\tilde{x}_a$ as under x_a . But the indifferent type ξ is higher. By $F(\xi)/f(\xi)$ and $F(\xi)$ being strictly increasing, equilibrium profits are strictly higher $\tilde{x}_a$ than under x_a .

We now know that, in equilibrium, firms choose product locations on the line of consumer's revealed preferences. We also know from Theorem 1 that the line of consumer preferences has direction $\Sigma A(x_b - x_a)$. If product locations x_a and x_b are on said line, then $x_b - x_a$ is parallel to $\Sigma A(x_b - x_a)$. This is the case if and only if $x_b - x_a$ is an eigenvector of ΣA . $\square$

C.2 Theorem 7

Proof. Fix a v_i , $i \in \{1, \dots, n\}$. We look for equilibria where product locations (x_a, x_b) are on the line spanned by v_i and consumer learning is to be about the A -projection

$$\frac{\langle v_i, \theta \rangle_A}{\langle v_i, v_i \rangle_A} \cdot v_i = \langle v_i, \theta \rangle_A \cdot v_i$$

of their ideal points on v_i . By Lemma 2 all equilibria are of this form.

To find such equilibria, we can reduce the model to one dimension by projecting the product

space on v_i . The utility of consumers for product location $\hat{x} \cdot v_i$ given ideal point $\hat{\theta} \cdot v_i$ is

$$-(\hat{x} \cdot v_i - \hat{\theta} \cdot v_i)^\top A(\hat{x} \cdot v_i - \hat{\theta} \cdot v_i) = -(\hat{x} - \hat{\theta})^2.$$

The distribution of ideal points projected on v_i is $\mathcal{N}(0, \sigma_\mu^2)$ with

$$\sigma_\mu^2 = v_i^\top A \Sigma A v_i.$$

Normal signal structures induce normal distributions $\rho = \mathcal{N}(0, \sigma_\rho^2)$ over posteriors means with posterior variance σ_π^2 , where $\sigma_\mu^2 = \sigma_\rho^2 + \sigma_\pi^2$ by the law of total variance.

Anderson, Goeree, and Ramer (1997) analyze a one-dimensional Hotelling model under quadratic consumer preferences $u(x, \theta) = -(x - \theta)^2$, $x, \theta \in \mathbb{R}$, and an exogenous log-concave distribution of consumer preferences. Using their Corollary 1, we know that given the distribution over revealed preferences $\rho = \mathcal{N}(0, \sigma_\rho^2)$ with density f_ρ , the unique equilibrium product attributes x and prices p are characterized by

$$x := -x_a = x_b = \frac{3}{4f_\rho(0)} = \frac{3}{4}\sqrt{2\pi}\sigma_\rho, \quad (39)$$

$$p := p_a = p_b = c \frac{3}{2f_\rho(0)^2} = 3\pi\sigma_\rho^2. \quad (40)$$

To solve for the optimal standard deviation of consumer preferences σ_ρ given x_a and x_b , we write the consumer's instrumental value of information of the normal distribution τ over normal posteriors π as as function of σ_ρ :

$$\begin{aligned} \mathbb{E}_\tau \left[\mathbb{E}_\pi \left[\max_{j \in \{a, b\}} \{-(\theta - x_j)^2\} - p \right] \right] &= \mathbb{E}_\tau \left[-(|\mathbb{E}_\pi[\theta]| - x)^2 - \sigma_\pi^2 - p \right] \\ &= \mathbb{E}_\tau \left[-\mathbb{E}_\pi[\theta]^2 + 2x\mathbb{E}_\pi[|\theta|] - x^2 - \sigma_\pi^2 - p \right] \\ &= -\sigma_\rho^2 + \sqrt{\frac{8}{\pi}}x\sigma_\rho - x^2 - \sigma_\pi^2 - p \\ &= \sqrt{\frac{8}{\pi}}x\sigma_\rho - x^2 - p - \sigma_\mu^2 \end{aligned}$$

The information cost is

$$\kappa (\log(2\pi\sigma_\mu^2) - \log(2\pi\sigma_\pi^2)) = \kappa (\log(\sigma_\mu^2) - \log(\sigma_\pi^2)).$$

Neglecting constants, the consumer chooses σ_ρ to solve

$$\begin{aligned} \max_{\sigma_\rho} \quad & \sqrt{\frac{8}{\pi}}x\sigma_\rho + \kappa \log(\sigma_\mu^2 - \sigma_\rho^2) \\ \text{s.t.} \quad & 0 \leq \sigma_\rho \leq \sigma_\mu. \end{aligned}$$

The second derivative of the objective

$$\frac{-2(\sigma_\mu^2 - \sigma_\rho^2) - 4\sigma_\rho^2}{(\sigma_\mu^2 - \sigma_\rho^2)^2} = \frac{-2\sigma_\mu^2 - 4\sigma_\rho^2}{(\sigma_\mu^2 - \sigma_\rho^2)^2} < 0,$$

is negative, so the first-order condition is sufficient for optimality. Supposing the inequality constraints are not binding, we get the first-order condition

$$\sqrt{\frac{8}{\pi}} \frac{x}{\kappa} = \frac{2\sigma_\rho}{\sigma_\mu^2 - \sigma_\rho^2}. \quad (41)$$

Defining $d := \sqrt{8/\pi}x/\kappa$, the equation has a unique positive solution

$$\sigma_\rho = -\frac{1}{d} + \sqrt{\frac{1}{d^2} + \sigma_\mu^2}$$

that lies between 0 and σ_μ , so the inequality constraints are satisfied. Thus, the first-order condition (41) characterizes the unique optimum of consumer learning.

Inserting equilibrium product locations (39) into the first-order condition (41) of consumer learning, we obtain

$$\frac{3}{\kappa}\sigma_\rho = \frac{2\sigma_\rho}{\sigma_\mu^2 - \sigma_\rho^2} \implies \sigma_\rho = 0 \quad \vee \quad \sigma_\mu^2 - \sigma_\rho^2 = \frac{2}{3}\kappa.$$

Thus, under $2\kappa/3 \geq \sigma_\mu^2 = v_i^\top A \Sigma A v_i$, there is a unique equilibrium without learning, product differentiation, and markups,

$$\sigma_\rho = 0, \quad x_a = x_b = 0, \quad p_a = p_b = 0.$$

Else, there is an additional equilibrium with learning, product differentiation, and markups,

$$\sigma_\rho^2 = \sigma_\mu^2 - \frac{2}{3}\kappa, \quad -x_a = x_b = \frac{3}{4}\sqrt{2\pi}\sigma_\rho, \quad p_a = p_b = 3\pi\sigma_\rho^2.$$

This concludes the proof. $\square$

C.3 Corollary 2

Proof. By the proof of Theorem 7, consumer utility (whether ex ante or ex post) is

$$U_i = -\sigma_\mu^2 + \sqrt{\frac{8}{\pi}}x\sigma_\rho - x^2 - p - \kappa(\log(\sigma_\mu^2) - \log(\sigma_\mu^2 - \sigma_\rho^2)).$$

When $\frac{2}{3}\kappa > \sigma_\mu^2$, then a small change in κ does not change the no-learning equilibrium and consumer utility remains $-\sigma_\mu^2$.

When $\frac{2}{3}\kappa \leq \sigma_\mu^2$, a change in κ has a direct effect on consumer welfare through the information cost and an indirect effect through σ_ρ , x , and p . By the envelope theorem, the indirect effect

through σ_ρ is zero. Using

$$x = \frac{3}{4}\sqrt{2\pi}\sigma_\rho = \frac{3}{4}\sqrt{2\pi}\sqrt{\sigma_\mu^2 - \frac{2}{3}\kappa},$$

$$p = 3\pi\sigma_\rho^2 = 3\pi\left(\sigma_\mu^2 - \frac{2}{3}\kappa\right),$$

the total derivative of consumer utility with respect to κ is

$$\begin{aligned}\frac{d}{d\kappa}U_i &= \frac{dx}{d\kappa}\frac{d}{dx}(-x^2) + \frac{dp}{d\kappa}\frac{d}{dp}(-p) + \frac{\partial}{\partial\kappa}(-\kappa(\log(\sigma_\mu^2) - \log(\sigma_\mu^2 - \sigma_\rho^2))) \\ &= \frac{dx}{d\kappa}(-2x) + \frac{dp}{d\kappa}(-1) - \left(\log(\sigma_\mu^2) - \log\left(\frac{2}{3}\kappa\right)\right) \\ &= -\frac{\sqrt{2\pi}}{8\sigma_\rho}\left(-\frac{3}{2}\sqrt{2\pi}\sigma_\rho\right) + \pi - \left(\log(\sigma_\mu^2) - \log\left(\frac{2}{3}\kappa\right)\right) \\ &= \left(1 + \frac{3}{8}\right)\pi - \left(\log(\sigma_\mu^2) - \log\left(\frac{2}{3}\kappa\right)\right)\end{aligned}$$

Thus, increasing the cost parameter κ has constant positive marginal effect of $(1+3/8)\pi$ on consumer utility through lowering product differentiation and prices, and a negative effect through its direct effect on the cost of information. The latter effect is becoming less strong as κ increases. Thus, voter utility is quasi-convex in κ .

Setting the total derivative to zero, we obtain that minimal consumer utility is achieved at

$$\left(1 + \frac{3}{8}\right)\pi = \log(\sigma_\mu^2) - \log\left(\frac{\kappa}{3}\right) \Rightarrow e^{(1+\frac{3}{8})\pi} = \frac{3\sigma_\mu^2}{\kappa} \Rightarrow \kappa = 3e^{-(1+\frac{3}{8})\pi}\sigma_\mu^2 \approx 0.04\sigma_\mu^2.$$

Thus, consumer utility is maximal under $\kappa \geq \frac{2}{3}\sigma_\mu^2$ or under $\kappa = 0$. Under $\kappa = \frac{2}{3}\sigma_\mu^2$, consumer utility is $-\sigma_\mu^2$. Under $\kappa = 0$, we have $\sigma_\rho = \sigma_\mu$ and consumer utility is

$$\begin{aligned}U_i &= -\sigma_\mu^2 + \sqrt{\frac{8}{\pi}}x\sigma_\rho - x^2 - p \\ &= -\sigma_\mu^2 + \sqrt{\frac{8}{\pi}}\frac{3}{4}\sqrt{2\pi}\sigma_\mu^2 - \frac{9}{8}\pi\sigma_\mu^2 - 3\pi\sigma_\mu^2 \\ &= \sigma_\mu^2\left(-1 + 3 - \frac{9}{8}\pi - 3\pi\right) < -\sigma_\mu^2.\end{aligned}$$

Thus, consumer welfare is maximal under $\kappa \geq \frac{2}{3}\sigma_\mu^2$. □

D Appendix: Other Results and Proofs

Throughout, we use the notation $\langle x, y \rangle_A := x^\top A y$.

D.1 Imperfect Observation of Party Platforms

Our results on voter learning are, under some assumptions, robust to voters observing a *noisy* signal about party platforms before learning about ideal points (under a timing where parties move before voters learn).

Suppose that platforms are stochastic and independent of each other and of voters' ideal points. This stochasticity may stem from parties making random errors when choosing their platforms or from parties having stochastic and private ideal points, similar to Matějka and Tabellini (2021). Further, assume that, after a common signal about platforms, voters acquire a signal about their ideal points. Because both signals preserve the independence of platforms and ideal points, the expected policy utility from voting for unknown platform x under unknown ideal point θ can be written as

$$\mathbb{E}[u(x, \theta)] = u(\mathbb{E}[x], \mathbb{E}[\theta]) - \mathbb{E}[(x - \mathbb{E}[x])^\top A(x - \mathbb{E}[x])] - \mathbb{E}[(\theta - \mathbb{E}[\theta])^\top A(\theta - \mathbb{E}[\theta])].$$

Up to the constant $\mathbb{E}[(x - \mathbb{E}[x])^\top A(x - \mathbb{E}[x])]$, the agent's utility is as under known platforms, (3), except for replacing the platform x with the expected platform $\mathbb{E}[x]$. Thus, our results from section 3 remain to hold after replacing party platforms with their expectation.

If voters observe heterogeneous signals about platforms, this creates heterogeneity in their learning strategies, such as the direction in which they learn (cf. Theorem 1). However, as long as their signals about platforms are similar enough, our results should carry over approximately. As a consequence, an extension of our model to heterogeneous signal may explain the empirical finding that the ideal points of politically better informed citizens are better described by a low-dimensional model (Converse, 1964; Hare, 2022), since better informed voters should have more homogeneous information about platforms.

D.2 Existence and Continuity

We show the set of maximizers of the voter learning problem are nonempty and upper hemicontinuous in the appropriate topology. We cannot show this using Berge's maximum theorem because, for an infinite state space, the Kullback-Leibler divergence, which is part of the objective, is only lower semicontinuous and not continuous. Berge's maximum theorem requires a continuous (and not just upper semicontinuous) objective to show upper hemicontinuity of the argmax correspondence. Although our objective is only upper semicontinuous in the choice variable, it is continuous in the parameter (the value function). Using this observation, we can apply the generalization of Berge's maximum theorem by Tian and Zhou (1992) to obtain our result.

To apply the result by Tian and Zhou (1992), we first note that the set of Bayes-consistent

distributions

$$X := \{\tau \in \Delta(\Delta(\mathbb{R}^n)) \mid \mathbb{E}_\tau[\pi] = \mu\}$$

is compact with respect to the weak topology by Kartik, Lee, Liu, and Rappoport (2022), Lemma 4. They show this result for any sigma-compact Polish state space by applying Prokhorov's theorem twice. For the interested reader, we include a shorter proof for the state space $\mathbb{R}^n$ by using a compactification argument.

Lemma 22. *The set of Bayes-consistent distributions X is compact with respect to the weak topology.*

Proof. Denote by $\mathbb{R}^n \cup \{\infty\}$ the one-point compactification of $\mathbb{R}^n$, which is an embedding. The space $\mathbb{R}^n \cup \{\infty\}$ is homeomorphic to the unit n -sphere, so it is a (compact) Polish space. By Aliprantis and Border (2006), Theorem 15.14, the pushforward of the embedding induces an embedding between Polish spaces $\Delta(\mathbb{R}^n) \hookrightarrow \Delta(\mathbb{R}^n \cup \{\infty\})$. By iteration of this argument, this induces an embedding $\Delta(\Delta(\mathbb{R}^n)) \hookrightarrow \Delta(\Delta(\mathbb{R}^n \cup \{\infty\}))$. Under this embedding, the image of X is the set $\hat{X} := \{\tau \in \Delta(\Delta(\mathbb{R}^n \cup \{\infty\})) \mid \int \pi d\tau = \mu\}$. The space $\Delta(\Delta(\mathbb{R}^n \cup \{\infty\}))$ is compact because $\mathbb{R}^n \cup \{\infty\}$ is a compact Polish space. The set $\hat{X}$ is the preimage of a singleton $\{\mu\}$ under a continuous function $\tau \mapsto \int \pi d\tau$, so it is closed. A closed subset of a compact space is compact, so $\hat{X}$ is compact. The set X is the preimage of $\hat{X}$ under an embedding, so X is compact. $\square$

Next, we prove a general maximum theorem for information-design problems on non-compact state spaces and for upper semicontinuous and bounded-from-above value functions. As above, define X as the set of Bayes-consistent distributions $\tau \in \Delta(\Delta(\mathbb{R}^n))$ over posteriors $\pi \in \Delta(\mathbb{R}^n)$ endowed with the topology induced by weak convergence. Define Y as the set of upper semicontinuous and bounded-from-above value functions v from $\Delta(\mathbb{R}^n)$ to $\mathbb{R} \cup \{-\infty\}$, endowed with the topology induced by uniform convergence. Here, upper semicontinuity is defined with respect to the topology of weak convergence on $\Delta(\mathbb{R}^n)$. Define $f: X \times Y \rightarrow \mathbb{R}$ as the expected value

$$f(\tau, v) = \int v(\pi) d\tau(\pi).$$

Proposition 7. *The argmax correspondence of the information design problem,*

$$\begin{aligned} M: Y &\rightrightarrows X \\ M(v) &:= \arg \max_{\tau \in X} f(\tau, v), \end{aligned}$$

is nonempty compact-valued and upper hemicontinuous.

Proof. For the proof, we use Theorem 1 in Tian and Zhou (1992). It shows that in a maximization problem, if (1) the objective is upper semicontinuous and feasible path transfer lower semicontinuous and (2) the feasibility correspondence is nonempty compact-valued, closed and upper hemicontinuous, then the maximum correspondence is nonempty compact-valued and upper hemicontinuous.

First, we show upper semicontinuity of f in (τ, v) . Suppose τ_n converges weakly to τ and v_n converges uniformly to v . We abbreviate $\int v(\pi)d\tau(\pi)$ by $\int v d\tau$ and show that $\lim_{n \rightarrow \infty} f(\tau_n, v_n) - f(\tau, v) \leq 0$:

$$\begin{aligned} \lim_{n \rightarrow \infty} \int v_n d\tau_n - \int v d\tau &= \lim_{n \rightarrow \infty} \left(\int v_n d\tau_n - \int v d\tau_n \right) + \left(\int v d\tau_n - \int v d\tau \right) \\ &= \lim_{n \rightarrow \infty} \int (v_n - v) d\tau_n + \left(\int v d\tau_n - \int v d\tau \right) \\ &\leq \lim_{n \rightarrow \infty} \int |v_n - v| d\tau_n + \lim_{n \rightarrow \infty} \left(\int v d\tau_n - \int v d\tau \right) \leq 0 \end{aligned} \quad (42)$$

In the last line, the first limit is zero because v_n converges uniformly to v and τ_n is a probability measure. By Villani (2009), Lemma 4.3, v being upper semicontinuous and bounded from above implies that $\tau \mapsto \int v d\tau$ is upper semicontinuous, so the second limit is less or equal to zero.

Second, we show feasible path transfer lower semicontinuity, which is introduced by Tian and Zhou (1992). The objective f is feasible path transfer lower semicontinuous in y with respect to feasibility correspondence F if for each $(x, y) \in X \times Y$ with $x \in F(y)$, there exists some neighborhood $\mathcal{N}(y)$ of y such that $\forall y' \in \mathcal{N}(y), \exists x' \in F(y')$ satisfying

$$f(x, y) \leq \liminf_{y' \rightarrow y} f(x', y').$$

Because in our case, $\forall y \in Y: F(y) = X$, we can choose $x' = x$ for all y' . Then, for any sequence $y_n \rightarrow y$

$$\lim_{n \rightarrow \infty} f(x, y_n) = \int y_n dx = \int y dx = f(x, y)$$

because y_n converges uniformly to y and x is a probability measure. So, $f(x, y) \leq \liminf_{y' \rightarrow y} f(x', y')$. Key is that while our objective may be discontinuous in the choice variable, it is continuous in the parameter (the value function).

Finally, in our case, the feasibility correspondence is nonempty and constant. Hence, it is closed and continuous, and thus also upper hemicontinuous. It is compact-valued by Lemma 22. $\square$

To apply Proposition 7, we show the value function of the voter learning problem,

$$\mathbb{E}_\nu \left[\max \left\{ \mathbb{E}_\pi[u(x_a, \theta)], \mathbb{E}_\pi[u(x_b, \theta)] + \nu \right\} \right] - \kappa D(\pi || \mu) \quad (43)$$

is indeed upper semicontinuous and bounded from above. The only property of the information cost that this result uses is lower semicontinuity of the Kullback-Leibler divergence.

Lemma 23. *The value function, (43), is upper semicontinuous in π and bounded from above.*

Proof. We separately show upper semicontinuity and boundedness from above for the information cost $-\kappa D(\pi || \mu)$ and for the instrumental value, $\mathbb{E}_\nu[\max\{\mathbb{E}_\pi[u(x_a, \theta)], \mathbb{E}_\pi[u(x_b, \theta)] + \nu\}]$, of information. Then, the their sum is upper semicontinuous and bounded from above.

The divergence $D_{\text{KL}}(\cdot||\mu)$ is bounded from below as it is non-negative and lower semicontinuous by Posner (1975), Theorem 1. Thus, $-D_{\text{KL}}(\pi||\mu)$ is upper semicontinuous and bounded from above.

Define the instrumental value of a posterior $V(\pi)$ as

$$V(\pi) = \mathbb{E}_\nu \left[\max \left\{ \mathbb{E}_\pi[u(x_a, \theta)], \mathbb{E}_\pi[u(x_b, \theta)] + \nu \right\} \right].$$

The utility $u(x, \theta) = -(x - \theta)^\top A(x - \theta)$ is continuous (and hence upper semicontinuous) in θ and bounded from above by zero. By Villani (2009), Lemma 4.3, the function $\pi \mapsto \mathbb{E}_\pi[u(x, \theta)] = \int u(x, \theta) d\pi(\theta)$ is upper semicontinuous in π , and it is bounded from above by zero. Further,

$$\mathbb{E}_\nu[\max\{l, r + \nu\}] = lF_\nu(l - r) + r(1 - F_\nu(l - r)) + \int_{l-r}^{\infty} s dF_\nu(s).$$

Because valence ν has a continuous distribution, F_ν is differentiable. Thus, $\mathbb{E}_\nu[\max\{l, r + \nu\}]$ is continuous in l and r . Thus, $V(\pi)$ is upper semicontinuous.

Further,

$$lF_\nu(l - r) + r(1 - F_\nu(l - r)) + \int_{l-r}^{\infty} s dF_\nu(s) \leq \max\{l, r\} + \frac{1}{2} \int |s| dF_\nu(s),$$

so by boundedness of $\mathbb{E}_\pi[u(x, \theta)]$ from above and ν having a finite first absolute moment, the instrumental value $V(\pi)$ is bounded from above. $\square$

Corollary 5. *A solution to the voter learning problem (P) exists.*

Proof. Follows immediately from Proposition 7 and Lemma 23. $\square$

Finally, for our proof of Theorem 2, we show the following result.

Lemma 24. *The value function, (43), converges uniformly as the valence shock converges in mean to zero.*

Proof. Let $\pi \in \Delta(\mathbb{R}^n)$. The information cost $-\kappa D(\pi||\mu)$ does not depend on the valence shock and can thus be ignored. The other component of the value function is the instrumental value, which we write as $V(\pi, F_\nu)$ as a function of the posterior π and CDF of ν ,

$$V(\pi, F_\nu) = \int_{-\infty}^{\infty} \max\{\mathbb{E}_\pi[u(x_a, \theta)], \mathbb{E}_\pi[u(x_b, \theta)] + s\} dF_\nu(s).$$

Write $l := \mathbb{E}_\pi[u(x_a, \theta)]$ and $r := \mathbb{E}_\pi[u(x_b, \theta)]$, so $V(\pi, F_\nu) = \int_{-\infty}^{\infty} \max\{l, r + s\} dF_\nu(s)$. For ν degenerate at 0, that is $F_\nu(\nu) = \mathbb{1}_{\{\nu \geq 0\}}$, we have $V(\pi, \mathbb{1}_{\{\nu \geq 0\}}) = \max\{l, r\}$. We show $V(\pi, F_\nu)$ converges to $V(\pi, \mathbb{1}_{\{\nu \geq 0\}})$ as ν converges to zero in mean.

First, consider the case that $l \geq r$, so $V(\pi, \mathbb{1}_{\{\nu \geq 0\}}) = l$. Then,

$$\begin{aligned}
l &= \int_{-\infty}^{\infty} l dF_{\nu}(s) \leq \int_{-\infty}^{\infty} \max\{l, r + s\} dF_{\nu}(s) \\
&= V(\pi, F_{\nu}) = l + \int_{-\infty}^{\infty} \max\{0, r - l + s\} dF_{\nu}(s) \\
&\leq l + \int_0^{\infty} \max\{0, s\} dF_{\nu}(s) = l + \frac{1}{2} \int_{-\infty}^{\infty} |s| dF_{\nu}(s) \\
&\Rightarrow l \leq V(\pi, F_{\nu}) \leq l + \frac{1}{2} \int_{-\infty}^{\infty} |s| dF_{\nu}(s),
\end{aligned}$$

where we have used the symmetry of ν in the third line.

Second, consider the case that $r > l$, so $V(\pi, \mathbb{1}_{\{\nu \geq 0\}}) = r$. By symmetry of the density of ν ,

$$\begin{aligned}
r &= \int_{-\infty}^{\infty} r dF_{\nu}(s) = \int_{-\infty}^{\infty} (r + s) dF_{\nu}(s) \leq \int_{-\infty}^{\infty} \max\{l, r + s\} dF_{\nu}(s) \\
&= V(\pi, F_{\nu}) = r + \int_{-\infty}^{\infty} \max\{l - r, s\} dF_{\nu}(s) \\
&\leq r + \int_{-\infty}^{\infty} \max\{0, s\} dF_{\nu}(s) = r + \frac{1}{2} \int_{-\infty}^{\infty} |s| dF_{\nu}(s) \\
&\Rightarrow r \leq V(\pi, F_{\nu}) \leq r + \frac{1}{2} \int_{-\infty}^{\infty} |s| dF_{\nu}(s).
\end{aligned}$$

Together, we have

$$|V(\pi, F_{\nu}) - V(\pi, \mathbb{1}_{\{\nu \geq 0\}})| = |V(\pi, F_{\nu}) - \max\{l, r\}| \leq \frac{1}{2} \int_{-\infty}^{\infty} |s| dF_{\nu}(s).$$

Thus, $V(\pi, F_{\nu})$ converges uniformly to $V(\pi, \mathbb{1}_{\{\nu \geq 0\}})$ as ν converges in mean to zero. $\square$

D.3 Distance-Based Information Costs

The proof of Theorem 1 uses only Blackwell monotonicity (step 3), posterior separability (step 2), and a notion of reflection invariance (step 1) of the information cost. Thus, the result holds for all information costs that satisfy these conditions. More precisely, step 3 of the proof uses that the Kullback-Leibler divergence is strictly convex in its first argument for posteriors with finite divergence, to show a strict mean-preserving contraction in posterior space strictly lowers the information cost. Step 2 of the proof uses linearity of the information cost under mixing between distributions over posteriors, which follows from posterior separability of the information cost. Step 1 uses that the Kullback-Leibler divergence is invariant under the constructed reflection Ref , $D_{\text{KL}}(\text{Ref}(\pi) || \text{Ref}(\mu)) = D_{\text{KL}}(\pi || \mu)$. This holds for all *invariant* divergences (Amari, 2016), which remain unchanged under any transformation of the state space. In fact, for invariant divergences, we sketch a somewhat shorter proof of Theorem 1 in the proof of Lemma 25. Other notable examples of invariant divergences, besides the Kullback-Leibler divergence, are the Rényi divergences, which have a foundation based on Blackwell dominance under repeated observation (Mu, Pomatto, Strack,

and Tamuz, 2021).

Furthermore, the proof of Theorem 1 generalizes beyond invariant divergences to certain distance-based divergences. Information costs based on invariant divergences have been criticized because they imply that any two states are equally costly to distinguish. The literature has, inspired by evidence from perceptual experiments, proposed distance-based information costs that make it more costly to distinguish between closer states (Hébert and Woodford, 2021; Pomatto, Strack, and Tamuz, 2023). Recall that in step 1 of our proof we use that the divergence D satisfies

$$D(\text{Ref}(\pi) || \text{Ref}(\mu)) = D(\pi || \mu)$$

where Ref is a reflection across a line that preserves the prior μ . Because the prior μ is elliptical with covariance matrix Σ , the prior is preserved by members of the orthogonal group of inner product Σ^{-1} , which is $\{Q \in \mathbb{R}^{n \times n} | Q^\top \Sigma^{-1} Q = \Sigma^{-1}\}$. Reflections across a line are those members of the orthogonal group that deliver the identity function when applied twice and that have a line as the subset of the space that is invariant under the mapping. If the divergence D is invariant under all such reflections, then step 1 of our proof works for it. We proceed to show this condition is satisfied for certain distance-based divergences. The upshot will be that the information cost needs to be based on a distance that is compatible with inner product Σ^{-1} .

While there is no generally agreed upon definition of distance-based costs, we assume that if a posterior-separable information cost is based on distance $d: \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$, then it should be invariant under *isometries* of d . Recall that a bijection $g: \mathbb{R}^n \rightarrow \mathbb{R}^n$ is an isometry of $d: \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}_{\geq 0}$ if $\forall v, w \in \mathbb{R}^n: d(v, w) = d(g(v), g(w))$. Intuitively, if we relabel the states such that the distance d is preserved, the cost of information should not change. This should be seen as a minimal implication for an information cost to be based on distance d , which makes our results stronger than had we imposed a stronger requirement. Below we show that versions of recent proposals for distance-based information costs satisfy this condition.

Definition 3 (Distance-Compatible Information Cost). *A posterior-separable information cost $c(\tau) = \mathbb{E}_\tau[D(\pi || \mu)]$ on state space $\mathbb{R}^n$ is compatible with distance $d: \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}_{\geq 0}$ if for all isometries $g: \mathbb{R}^n \rightarrow \mathbb{R}^n$ of d*

$$D(\pi || \mu) = D(g(\pi) || g(\mu)),$$

where $g: \Delta(\mathbb{R}^n) \rightarrow \Delta(\mathbb{R}^n)$ is the pushforward induced by g .

Particularly important in our context are the standard Euclidean distance on $\mathbb{R}^n$, $d(v, w) = ((v - w)^\top (v - w))^{1/2}$, and the (non-standard) Euclidean distance induced by a symmetric, positive definite matrix C , $d(v, w) = (v - w)^\top C (v - w)$. That is because we need the information cost to be compatible with a metric d of the state space $\mathbb{R}^n$ such that the reflection Ref in our proof is an isometry of d . Using the fact that Ref is an isometry of the distance induced by Σ^{-1} , we can easily prove the following corollary.

Corollary 6. *Let the information cost be posterior separable, $c(\tau) = \mathbb{E}_\tau[D(\pi || \mu)]$, consistent with the standard Euclidean distance, the divergence D be strictly convex in its first argument, and the*

prior be spherical, $\Sigma = I_n$. Then, the induced posterior means are on the line through the prior mean and $A(x_b - x_a)$ under any optimal information τ .

More generally, maintaining posterior separability, and convexity of the divergence D , let the information cost be consistent with the Euclidean distance induced by symmetric, positive definite matrix C and the prior be spherical in an orthonormal basis of C , that is Σ^{-1} is a multiple of C . Then, the posterior means are on a line through the prior mean and $\Sigma A(x_b - x_a)$ under any optimal information τ .

Proof. Ref is a reflection with respect to Σ^{-1} and thus an isometry of the distance induced by Σ^{-1} . By $\Sigma^{-1} = kC$ with $k \in \mathbb{R}$, Ref is also an isometry of the distance induced by C . Thus, $D(\pi||\mu) = D(\text{Ref}(\pi)||\text{Ref}(\mu))$ and our proof of Theorem 1 applies. $\square$

Intuitively, our proof of Theorem 1 requires that a reflection that preserves the prior (so the reflection defines a Bayes-consistent distribution over posteriors) to preserve the information cost. The elliptical prior with covariance matrix Σ is preserved by reflections with respect to the inner product induced by Σ^{-1} . Thus, we need the information-cost distance to induce the same geometry as Σ^{-1} . This is the case if the information-cost distance is based on an inner product induced by matrix C where C is a multiple of Σ^{-1} .

Finally, we present a few examples of information costs that are compatible with the Euclidean distance induced by a symmetric, positive definite matrix C . Strict convexity of these divergences, for posteriors with finite divergence, is either known or can be shown. (The first example is strictly convex only for posteriors that do not share the same mean, which suffices for the proof of Theorem 1.)

Example 1: Posterior-Variance Cost A multidimensional version of the posterior-variance cost can be defined as the divergence D_{Var} being a second central moment (a generalization of variance to arbitrary metric spaces),

$$D_{\text{Var}}(\pi||\mu) = \mathbb{E}_{\pi} \left[-d(\theta, \mathbb{E}_{\pi}[\theta])^2 \right] = \int -d(\theta, \mathbb{E}_{\pi}[\theta])^2 d\pi(\theta),$$

where $d(v, w) = \sqrt{v^{\top} C w}$ and C is a symmetric, positive definite $n \times n$ -matrix. Under any bijection g on $\mathbb{R}^n$ that preserves the inner product $\langle \cdot, \cdot \rangle_C$, the information cost is preserved, that is $D(g(\pi)||g(\mu)) = D(\pi||\mu)$, which can be seen by

$$\begin{aligned} D_{\text{Var}}(g(\pi)||g(\mu)) &= \int -d(\theta, \mathbb{E}_{\pi}[\theta])^2 dg(\pi)(\theta) \\ &= \int -d(g(\theta), \mathbb{E}_{\pi}[g(\theta)])^2 d\pi(\theta) \\ &= \int -d(g(\theta), g(\mathbb{E}_{\pi}[\theta]))^2 dg(\pi)(\theta) \\ &= \int -d(\theta, \mathbb{E}_{\pi}[\theta])^2 d\pi(\theta) = D_{\text{Var}}(\pi||\mu). \end{aligned}$$

Here, we have used that a bijection that preserves an inner product is linear and hence commutes with the expectation operator. That Examples 2 and 3 below are compatible with the Euclidean distance can shown in a similar fashion.

Example 2: Neighborhood-based Cost For state space $\mathbb{R}^n$, Hébert and Woodford (2021) propose the Fisher information cost, based on divergence

$$D(\pi||\mu) = \int_{\text{supp}(\pi)} c(\theta) \frac{|\nabla f_\pi(\theta)|^2}{f_\pi(\theta)} d\theta$$

for posteriors π with density f_π (and infinite divergence if the posterior is not absolutely continuous with respect to the Lebesgue measure) and where $c(\theta)$ captures how costly it is locally to differentiate between states. If $c(\theta) = c$ is constant, one can show similar to above that the cost is based the standard Euclidean distance.

Example 3: Log-Likelihood Ratio Cost The log-likelihood ratio (LLR) cost, introduced and axiomatized by Pomatto, Strack, and Tamuz (2023) is defined for finite state spaces Θ . Pomatto, Strack, and Tamuz (2023) show, given a full-support prior, the LLR cost is posterior-separable with divergence

$$D_{LLR}(\pi||\mu) = \sum_{\theta, \theta' \in \Theta} \beta(\theta, \theta') \frac{\pi(\theta)}{\mu(\theta)} \log \left(\frac{\pi(\theta)}{\pi(\theta')} \right),$$

if $\pi(\theta) > 0$ for all $\theta \in \Theta$ and infinite otherwise. The coefficients $\beta(\theta, \theta')$ capture how hard distinguishing between states θ and θ' is. If $\Theta \subset \mathbb{R}^n$ and $\beta(\theta, \theta') = f((\theta - \theta')^\top C(\theta - \theta'))$ for some function f , then it can be easily shown that the cost function is based on the Euclidean distance induced by C . While finite state space is not appropriate for our analysis since we assume the prior μ is elliptical, we conjecture that under an appropriate generalization to infinite state spaces, the resulting LLR cost remains distance-based in our sense.

D.4 Corollary 1

Proof. We denote the k party platforms by x_j , $j \in \{1, \dots, k\}$ and assume the utility of voting for candidate j under ideal point θ is $u(x_j, \theta) + \nu_j$. Define the valence vector $\nu := (\nu_j)_{j \in \{1, \dots, k\}}$, which can have an arbitrary distribution.

Analogously to the proof of Theorem 1, we first show the instrumental value of information τ depends only on the projection of the induced distribution over posterior means on the $(k - 1)$ -

dimensional subspace spanned by $\{x_j - x_1\}_{j=2,\dots,k}$.

$$\begin{aligned}
& \mathbb{E}_\tau \left[\mathbb{E}_\nu \left[\max_j \left\{ \mathbb{E}_\pi \left[-\langle \theta - x_j, \theta - x_j \rangle_A \right] + \nu_j \right\} \right] \right] \\
&= \mathbb{E}_\tau \left[\mathbb{E}_\nu \left[\max_j \left\{ -\langle x_j, x_j \rangle_A + 2\langle x_j, \mathbb{E}_\pi[\theta] \rangle_A + \nu_j \right\} \right] \right] - \mathbb{E}_\mu [\langle \theta, \theta \rangle_A] \\
&= \mathbb{E}_\tau \left[\mathbb{E}_\nu \left[\max_j \left\{ -\langle x_j, x_j \rangle_A + 2\langle x_j - x_1, \mathbb{E}_\pi[\theta] \rangle_A + \nu_j \right\} \right] + 2\langle x_1, \mathbb{E}_\pi[\theta] \rangle_A \right] - C_1 \\
&= \mathbb{E}_\tau \left[\mathbb{E}_\nu \left[\max_j \left\{ -\langle x_j, x_j \rangle_A + 2\langle x_j - x_1, \mathbb{E}_\pi[\theta] \rangle_A + \nu_j \right\} \right] \right] - C_1 + C_2
\end{aligned}$$

where $C_1 = \mathbb{E}_\mu [\langle \theta, \theta \rangle_A]$ and $C_2 = 2\langle x_1, \mathbb{E}_\mu[\theta] \rangle_A$ are constants and in fact zero under our prior. In the second line, we used the law of iterated expectations. In the last line, we used that the inner product is linear, to apply another law of iterated expectations. Thus, only the projection of $\mathbb{E}_\pi[\theta]$ on $\{x_j - x_1\}_{j=2,\dots,k}$ is payoff-relevant.

To replicate the second part of the proof of Theorem 1, we need to again define an appropriate reflection that preserves the instrumental value of information as well as the prior. Analogously to above, we let $\Delta \hat{x}_j := \Sigma A(x_j - x_1)$ and define the reflection as

$$\text{Ref}_k(\theta) = 2 \sum_{j=1}^k \frac{\langle \Delta \hat{x}_j, \theta \rangle_{\Sigma^{-1}}}{\langle \Delta \hat{x}_j, \Delta \hat{x}_j \rangle_{\Sigma^{-1}}} \Delta \hat{x}_j - \theta.$$

If $\Sigma = A = I_n$, this is just the standard reflection across the space spanned by $\{x_j - x_1\}_{j=2,\dots,k}$. In general, it is the suitable reflection that preserves the A -projection on this subspace as well as the prior (since it is a Σ^{-1} -reflection). The second part of the proof of Theorem 1 applies using the reflection Ref_k instead of Ref . $\square$

D.5 Proposition 1

Proof. First, we show the conclusion of Theorem 1 still holds under the restriction to normal distributions. This implies voters' candidates for optimal signal structures are one-dimensional and normal. Hence, the candidate signal structures are completely Blackwell-ordered and thereby ordered by information cost.

Lemma 25. *Restrict the prior μ and feasible signal structures to be normal. The conclusion of Theorem 1 still holds; that is, revealed voter ideal points are on the line through the prior mean with direction $\Sigma A(x_b - x_a)$.*

Proof. Under the restriction to normal signal structures, the reflection argument underlying our proof of Theorem 1 does not hold anymore because the better signal structure, constructed by that proof, need not be normal. Instead, we apply an argument based on a so-called pre-garbling, which shows for all invariant information costs such as mutual information that agents learn only about the partition of payoff-equivalent states (Amari, 2016; Caplin, Dean, and Leahy, 2022). If the acquired

signal structure was not measurable with respect to the partition of payoff-equivalent states, one can construct a better signal structure based on a pre-garbling, that is by, for each state, obtaining the average distribution over signals conditional on the partition element of the state (Caplin, Dean, and Leahy, 2022).³⁸ The resulting signal structure does not distinguish between payoff-equivalent states and it is better because it is cheaper and equally instrumentally valuable. It is not hard to see that such a pre-garbling maintains normality of the signal structure. Thus, also under the restriction to normal signal structures, the optimal signal structure is measurable with respect to the partition of payoff-equivalent states, which in our case are those states θ that have the same A -projection $\langle x_b - x_a, \theta \rangle_A$ on the platform difference $x_b - x_a$. To obtain the result of Theorem 1, two additional steps are necessary. First, suppose the voter learns the A -projection $S = \langle x_b - x_a, \theta \rangle_A$ of the state θ on $x_b - x_a$ perfectly. Upon learning $S = s$, by joint normality, the posterior mean would be

$$\mathbb{E}[\theta|S = s] = \mathbb{E}[\theta] + (s - \mathbb{E}[S]) \cdot \frac{\text{Cov}(\theta, S)}{\text{Var}(S)} = c(s) \text{Cov}(\theta, S) = c(s) \Sigma A(x_b - x_a),$$

where $c(s)$ is a scalar and we have used the normalization $\mathbb{E}[\theta] = 0$. Thus, the posterior means induced by S are on the line characterized by Theorem 1. Second, given that the voter actually acquires some garbling of S , the induced distribution over posterior means is a mean-preserving contraction of the one induced by S . Hence, the resulting distribution over posteriors means is also supported on the line characterized by Theorem 1. $\square$

Second, we show the comparative statics regarding the cost parameter κ . The voter's objective is supermodular in κ and in the cost of information $c(\tau)$. Thus, a smaller κ implies a greater cost of information $c(\tau)$ in the strong order. The one-dimensional normal distributions over posterior means are completely ordered by the mean-preserving spread relation, which coincides with the ordering by variance. Thus, a greater cost of information implies a greater variance.

Third, we show the comparative statics regarding the degree of platform polarization α . By Lemma 25, the distribution ρ of revealed ideology can be written as $\rho = X\mathcal{N}(0, \sigma_\rho^2)$ with $X := \frac{\Sigma X(x_b - x_a)}{\|\Sigma X(x_b - x_a)\|} = \frac{\Sigma X(x_b^* - x_a^*)}{\|\Sigma X(x_b^* - x_a^*)\|}$. We show the value of information is supermodular in the standard deviation of revealed ideology σ_ρ and the degree of platform polarization α . This implies that the optimal variance is increasing in the strong set order in α . Because the cost of information does not depend on α , it is sufficient to show the instrumental value of information is supermodular in σ_ρ and α . By $x_b^\top A x_b = x_a^\top A x_a$, we have $\langle x_b - x_a, \frac{x_a + x_b}{2} \rangle_A = 0$. Using this and (16), the instrumental

³⁸To be more specific, let S be a normal signal (modeled as a random vector), that is (S, θ) are jointly normal. Let $P(\theta) := \langle x_b - x_a, \theta \rangle_A = (A(x_b - x_a))^\top \theta$. The distribution of the pre-garbling $\tilde{S}$ conditional on some state θ should be identical to the distribution of S conditioned on the partition of payoff-equivalent states of θ , $\{\theta' \in \mathbb{R}^n | P(\theta') = P(\theta)\}$. Formally, we have

$$S|P(\theta) \sim \mathcal{N}(\mathbb{E}[S] + \Sigma_{SP}\Sigma_P^{-1}(P(\theta) - \mathbb{E}[P(\theta)]), \Sigma_S - \Sigma_{SP}\Sigma_P^{-1}\Sigma_{PS}).$$

where Σ_{SP} , Σ_P , and Σ_S are the relevant (cross-)covariance matrices of $P(\theta)$ and S . We can define the pre-garbling $\tilde{S}$ via

$$\tilde{S} = \mathbb{E}[S|P(\theta)] + \varepsilon$$

where $\varepsilon \sim \mathcal{N}(0, \Sigma_S - \Sigma_{SP}\Sigma_P^{-1}\Sigma_{PS})$ is independent of θ . Signal and state $(\tilde{S}, \theta)$ are jointly normal.

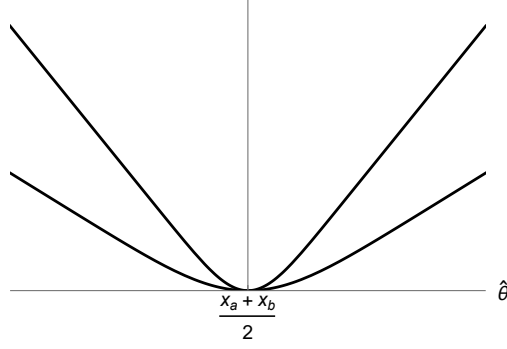


Figure 8: The value function $v(\hat{\theta}, \alpha)$ as a function of $\hat{\theta}$ for two values of α . The higher alpha corresponds to the steeper value function, which offers a higher return to information.

value of information as a function of σ_ρ and α , $V(\sigma_\rho, \alpha)$, can be written as an expectation of the one-dimensional variable $\hat{\theta}$,

$$\begin{aligned} V(\sigma_\rho, \alpha) &= \mathbb{E}_{\theta \sim X\mathcal{N}(0, \sigma_\rho^2)} \left[\mathbb{E}_\nu \left[\max \left\{ -\langle \alpha(x_b^* - x_a^*), \theta \rangle_A, \langle \alpha(x_b^* - x_a^*), \theta \rangle_A + \nu \right\} \right] \right] \\ &= \mathbb{E}_{\hat{\theta} \sim \mathcal{N}(0, \sigma_\rho^2)} \left[v(\hat{\theta}, \alpha) \right] = \mathbb{E}_{Z \sim \mathcal{N}(0, 1)} \left[v(\sigma_\rho Z, \alpha) \right], \end{aligned}$$

with

$$v(\hat{\theta}, \alpha) := \mathbb{E}_\nu \left[\max \left\{ -\alpha \underbrace{\langle x_b^* - x_a^*, X \rangle_A}_{=: \Delta} \hat{\theta}, \alpha \underbrace{\langle x_b^* - x_a^*, X \rangle_A}_{=: \Delta} \hat{\theta} + \nu \right\} \right].$$

For illustration, we graph the value function $v(\hat{\theta}, \alpha)$ as a function of $\hat{\theta}$ for different values of α in Figure 8. To show the instrumental value $V(\sigma_\rho, \alpha)$ is supermodular in σ_ρ and α , we show $\frac{d^2}{d\alpha d\sigma_\rho} V(\sigma_\rho, \alpha) > 0$. We have

$$\frac{d^2}{d\alpha d\sigma_\rho} V(\sigma_\rho, \alpha) = \mathbb{E}_{Z \sim \mathcal{N}(0, 1)} \left[Z \frac{d^2}{d\alpha d\hat{\theta}} v(\sigma_\rho Z, \alpha) \right]. \quad (44)$$

Assuming $\Delta > 0$ (the other case is analogous), the derivative of the value function $v(\hat{\theta}, \alpha)$ in $\hat{\theta}$ is, using the envelope theorem,

$$\frac{d}{d\hat{\theta}} v(\hat{\theta}, \alpha) = -\alpha\Delta + 2\alpha\Delta F_\nu(\alpha\Delta\hat{\theta}) = 2\alpha\Delta \left(F_\nu(\alpha\Delta\hat{\theta}) - \frac{1}{2} \right).$$

By symmetry of ν , the factor in brackets has the same sign as $\hat{\theta}$. Both factors are increasing in absolute value in α . Thus, this term is positive and increasing in α if $\hat{\theta} > 0$, zero if $\hat{\theta} = 0$, and negative and decreasing in α if $\hat{\theta} < 0$. Because (44) includes Z as an additional factor, the cross-derivative is positive. $\square$

Finally, we argue that our result is, under some conditions, robust to a common component of ideal point. Suppose there is both a common component and an idiosyncratic component of the ideal point, $\theta_i = \omega + \delta_i$ where $\omega \sim \mathcal{N}(0, \Sigma_\omega)$, $\delta_i \sim \mathcal{N}(0, \Sigma_\delta)$, and ω and all δ_i are mutually independent. By the proof of Lemma 25, voters acquire noisy signals about the A -projection of their ideal point on

the platform difference, $\hat{\theta} := (A(x_b - x_a))^\top \theta$. Defining $\hat{\omega} := (A(x_b - x_a))^\top \omega$ and $\hat{\delta} := (A(x_b - x_a))^\top \delta$ analogously, we can treat the voter's learning problem as a one-dimensional problem where they learn about $\hat{\theta} = \hat{\omega} + \hat{\delta}$ where $\hat{\omega} \sim \mathcal{N}(0, \sigma_\omega^2)$ and $\hat{\delta} \sim \mathcal{N}(0, \sigma_\delta^2)$, with $\sigma_\omega^2 = (A(x_b - x_a))^\top \Sigma_\omega A(x_b - x_a)$ and $\sigma_\delta^2 = (A(x_b - x_a))^\top \Sigma_\delta A(x_b - x_a)$. It remains to show the variance of the idiosyncratic uncertainty in (??) is increasing in the informativeness of the signal, or decreasing in the noise variance σ_ε^2 of the normal signal $S = \hat{\theta} + \varepsilon$. This is the case if and only if

$$\begin{aligned} \frac{d}{d\sigma_\varepsilon^2} \left(\frac{\sigma_\omega^2 + \sigma_\delta^2}{\sigma_\omega^2 + \sigma_\delta^2 + \sigma_\varepsilon^2} \right)^2 (\sigma_\delta^2 + \sigma_\varepsilon^2) < 0 &\Leftrightarrow \frac{d}{d\sigma_\varepsilon^2} \frac{\sigma_\delta^2 + \sigma_\varepsilon^2}{(\sigma_\omega^2 + \sigma_\delta^2 + \sigma_\varepsilon^2)^2} < 0 \\ &\Leftrightarrow (\sigma_\omega^2 + \sigma_\delta^2 + \sigma_\varepsilon^2)^2 - 2(\sigma_\delta^2 + \sigma_\varepsilon^2)(\sigma_\omega^2 + \sigma_\delta^2 + \sigma_\varepsilon^2) < 0 \Leftrightarrow \sigma_\omega^2 < \sigma_\delta^2 + \sigma_\varepsilon^2, \end{aligned}$$

which is implied by $\sigma_\delta^2 > \sigma_\omega^2$. That is, the monotone comparative statics holds as long as the variance σ_δ^2 of the idiosyncratic component (when A -projected on the platform difference $x_b - x_a$) is greater than the variance σ_ω^2 of the common component (when A -projected on the platform difference $x_b - x_a$).

D.6 Proposition 2

For brevity, we refer by property L to the property that ideal points are on a line when projected onto the space spanned by the survey questions, as defined in the main text.

It is known (e.g. Ladha, 1991) that the one-dimensional spatial model with quadratic utility is equivalent to the one-dimensional item-response theory (IRT) model, which the empirical papers we referred to estimate (Jessee, 2009; Jessee, 2012; Tausanovitch and Warshaw, 2012; Shor and Rogowski, 2018; Fowler, Hill, Lewis, Tausanovitch, Vavreck, and Warshaw, 2022).

The IRT model is given as follows. Let y_{ij} denote the response of individual i to question j , which can be either 1 or 0. Under a one-dimensional IRT model, the likelihood is given by

$$\Pr(y_{ij} = 1) = \Phi(\alpha_j + \beta_j p_i),$$

where Φ is the logistic or the normal cumulative distribution function, $\alpha_j, \beta_j \in \mathbb{R}$ are question-specific parameters, and $p_i \in \mathbb{R}$ are individual-specific parameters.

In our proof, we show (Part I) if a multidimensional spatial model has property L , then there is a one-dimensional IRT model that is observationally equivalent, and (Part II) if there is a one-dimensional IRT model that is observationally equivalent to a given multidimensional spatial model, then the spatial model satisfies property L . Because of the equivalent to of one-dimensional IRT models to one-dimensional spatial models with quadratic utility, this establishes our result.

Proof. First, in the spatial model, we have

$$\langle x_{j1} - \theta_i \rangle - \langle x_{j2} - \theta_i \rangle = \langle \Delta x_j, \theta_i - \bar{x}_j \rangle$$

where $\Delta x_j := x_{j1} - x_{j2}$ and $\bar{x}_j := \frac{x_{j1} + x_{j2}}{2}$. We suppose that the same response item is not part of

two different questions, which seems to be satisfied in practice, so there are no restrictions that x_{j1} or x_{j2} for different j are the same.³⁹ Hence, there are no restrictions on Δx_j and $\bar{x}_j$, and we can reparametrize the model through $\{\Delta x_j, \bar{x}_j, \theta_i\}$ instead of $\{x_{j1}, x_{j2}, \theta_i\}$.

The multidimensional spatial model with parameters $\{\Delta x_j, \bar{x}_j, \theta_i\}$ is observationally equivalent to the one-dimensional IRT model with parameters $\{\alpha_j, \beta_j, p_i\}$ if and only if

$$\forall i, j: \langle \Delta x_j, \theta_i - \bar{x}_j \rangle = \alpha_j + \beta_j p_i \quad (45)$$

Part I Suppose the multidimensional spatial model satisfies property L . Then, define

$$\begin{aligned} p_i &:= \lambda_i \\ \beta_j &:= \langle \Delta x_j, \Delta \theta \rangle \\ \alpha_j &:= \langle \Delta x_j, \theta_1 - \bar{x}_j \rangle \end{aligned}$$

Then, (45) holds by

$$\langle \Delta x_j, \theta_j - \bar{x}_j \rangle = \langle \Delta x_j, \theta_1 - \bar{x}_j + \lambda_i \Delta \theta + \theta_i^\perp \rangle = \alpha_j + \beta_j \lambda_i$$

and the IRT model with parameters $\{\alpha_j, \beta_j, p_i\}$ is observationally equivalent.

Part II If there is a one-dimensional model that is observationally equivalent to the multidimensional spatial model, then (45) holds, which implies

$$\langle \Delta x_j, \theta_i - \theta_1 \rangle = \beta_j (p_i - p_1).$$

Take $i = 2$, then the projection of $\theta_2 - \theta_1$ on all Δx_j is given, so $\theta_2 - \theta_1$ is uniquely pinned down in the space spanned by $\{\Delta x_j\}$. For any other $i > 2$, $\theta_i - \theta_1$ is $\beta_j (p_i - p_1) = \frac{p_i - p_1}{p_2 - p_1} \langle \Delta x_j, \theta_i - \theta_1 \rangle$. Thus, the projection of $\theta_i - \theta_1$ on the space spanned by $\{\Delta x_j\}$ is a multiple of the one of $\theta_2 - \theta_1$. Thus, property L holds. $\square$

D.7 Sufficiency of First-Order Conditions for Equilibrium Platforms

Recall that in the context of Theorem 3, all our *equilibrium candidates*, that is pairs of platforms (x_a, x_b) that satisfy the necessary first-order conditions of optimality, are of the form

$$(x_a, x_b) = \alpha(x_a^*, x_b^*)$$

with $\alpha \in [0, 1]$. It remains to show these platforms are indeed best responses to each other, that is, the first-order conditions are sufficient for optimality in these cases.

First, we show in Lemma 26 that it is sufficient for equilibrium candidates to be equilibria that party objectives are quasi-concave on certain compact subsets of $\mathbb{R}^n$. Second, we show in Lemma 27 that when the weight on vote share m is small enough or the valence shock ν is large enough,

³⁹This condition is sufficient but not necessary for our proof.

then this condition for quasi-concavity is satisfied. By comparison to existing results (Lindbeck and Weibull, 1987; Enelow and Hinich, 1989), our proofs are complicated by the fact that voter ideal points are not bounded because we assume a normal distribution in Theorem 3.

For the first part, we make use of the fact that any platform choice x_j outside the ellipse defined by $u(x_j, x_j^*) \geq -m$ is suboptimal, as we have shown at the beginning of the proof of Lemma 1. To formulate the sufficient condition for equilibrium candidates to be equilibria, define

$$\begin{aligned}\mathcal{E}_a &:= \{x \in \mathbb{R}^n | (x - x_a^*)^\top A(x - x_a^*) \leq m\} \\ \mathcal{E}_b &:= \{x \in \mathbb{R}^n | (x - x_b^*)^\top A(x - x_b^*) \leq m\}.\end{aligned}$$

Further, recall that because of the normal prior $\mu = \mathcal{N}(0, \Sigma)$ and the restriction to normal signals, the distribution ρ of posterior means is necessarily normal. Also, it is supported on the line through the origin with direction $\Sigma A(x_b - x_a)$. By the law of total variance, the variance of the normal distribution is bounded by the variance of the prior in that direction. Let R the set of all distributions ρ satisfying these three requirements.

With these definitions, we can state the sufficient (but not necessary) condition for equilibrium candidates to be equilibria.

Condition 1. *The following holds.*

- $U_a(x_a, x_b, \rho)$ is quasi-concave in x_a on $\mathcal{E}_a$ for all $x_b = \alpha x_b^*$ with $\alpha \in [0, 1]$ and $\rho \in R$.
- $U_b(x_a, x_b, \rho)$ is quasi-concave in x_b on $\mathcal{E}_b$ for all $x_a = \alpha x_a^*$ with $\alpha \in [0, 1]$ and $\rho \in R$.

Lemma 26. *If Condition 1 holds, then all equilibrium candidates are equilibria.*

Proof. For an equilibrium candidate, $(x_a, x_b) = \alpha(x_a^*, x_b^*)$ with $\alpha \in [0, 1]$, party a 's best response is in the ellipse $\mathcal{E}_a$, as argued above. Quasi-concave utility over $\mathcal{E}_a$ under the equilibrium x_b and any feasible ρ implies that the first-order condition is sufficient for optimality of x_a . The same holds for x_b . Thus, the equilibrium candidate (x_a, x_b, ρ) is an equilibrium as both platforms are best responses. $\square$

Next, we give assumptions that ensure that Condition 1 holds. Lemma 27 uses the following assumption on the density of the valence shock, which is satisfied for example by the normal density and the Laplace or double exponential density. This assumption is far from necessary but it suffices to show certain terms vanish faster than a polynomial term diverges, which we use in our proof.

Assumption 1. *The density of the valence shock $f_\nu(x)$ is proportional to $\exp\{-g(x)\}$ where $g(x)$ is of the form*

$$g(x) = c_0 - c_1|x| - c_2x^2 - c_3|x|^3 - \dots - c_m \cdot \begin{cases} |x|^m & \text{if } m \text{ odd} \\ x^m & \text{if } m \text{ even} \end{cases}$$

with $c_1, \dots, c_m \geq 0$.

In this definition we take the absolute value of the odd polynomial terms, as present for example in the Laplace density, to ensure that the valence shock is symmetric. We assume that the constants c_1 through c_m are positive to ensure that the density is quasi-concave.

Lemma 27. *Under Assumption 1, if the weight on vote share m is small enough or if the valence shock ν is large enough, then Condition 1 holds. Formally, there exist $\underline{m} > 0$, such that for all m with $0 < m < \underline{m}$, Condition 1 holds. Given a valence shock ν that satisfies our assumptions, there exist $K > 0$, such that for all $k > K$, Condition 1 holds under valence shock $k\nu$.*

Proof. We show the Hessian of the party objective $U_a(x_a, x_b, \rho)$ in x_a is negative definite over $\mathcal{E}_a$ when $x_b = \alpha x_b^*$ with $\alpha \in [0, 1]$ and $\rho = \mathcal{N}(0, \Sigma_\rho)$ with $\Sigma_\rho \leq \Sigma$. This implies concavity and therefore quasi-concavity. The proof for concavity of $U_b(x_a, x_b, \rho)$ in x_b is analogous by symmetry.

Let ∇ denote the gradient with respect to x_a . The Hessian $H_a(x_a, x_b, \rho, m)$ of party a 's objective with respect to x_a is

$$\begin{aligned} H_a(x_a, x_b, \rho, m) &:= \nabla^2 U_a(x_a, x_b, \rho) = m \int \nabla^2 F_\nu(\Delta u(\theta, x_a, x_b)) d\rho(\theta) + \nabla^2 u(x_a, x_a^*) \\ &= m \int \nabla (\nabla u(x_a, \theta) f_\nu(\Delta u(\theta, x_a, x_b))) d\rho(\theta) - I_n \\ &= m \int \left(-I_n f_\nu(\Delta u(\theta, x_a, x_b)) + \nabla u(x_a, \theta) \nabla u(x_a, \theta)^\top f'_\nu(\Delta u(\theta, x_a, x_b)) \right) d\rho(\theta) - I_n \\ &= m \int \left(-I_n f_\nu(\Delta u(\theta, x_a, x_b)) + 4(\theta - x_a)(\theta - x_a)^\top f'_\nu(\Delta u(\theta, x_a, x_b)) \right) d\rho(\theta) - I_n \end{aligned}$$

where

$$\Delta u(\theta, x_a, x_b) := u(x_a, \theta) - u(x_b, \theta).$$

We were able to exchange integration and differentiation because the derivative and second derivative are bounded (componentwise) by a constant, which is integrable under the probability measure ρ . We can bound both derivatives by constants because they consist of polynomial terms multiplied by an exponential function (with a decreasing polynomial exponent), so the integrands are eventually radially decreasing by Assumption 1. Thus, the supremum of the integrand is obtained on a compact sphere, on which the integrand obtains its finite maximum by continuity.

If $m = 0$, the Hessian $H_a(x_a, x_b, \rho, m)$ is $-A$, which is negative definite, so the objective is concave. To show for m small enough, the objective is concave, we first prove that the Hessian $H_a(x_a, x_b, \rho, m)$ is continuous in (x_a, x_b, ρ, m) . For that we endow the domain of ρ , $\Delta(\mathbb{R}^n)$, with the weak topology on $\Delta(\mathbb{R}^n)$.

First, we show continuity of the Hessian in (x_a, x_b, ρ) . Let the sequence $(x_a^n, x_b^n, \rho^n, m)_{n \in \mathbb{N}}$ converge to (x_a, x_b, ρ, m) . We show that

$$\begin{aligned} &\lim_{n \rightarrow \infty} (H_a(x_a, x_b, \rho, m) - H_a(x_a^n, x_b^n, \rho^n, m)) \\ &= \lim_{n \rightarrow \infty} (H_a(x_a, x_b, \rho, m) - H_a(x_a, x_b, \rho^n, m)) + \lim_{n \rightarrow \infty} (H_a(x_a, x_b, \rho^n, m) - H_a(x_a^n, x_b^n, \rho^n, m)) \quad (46) \\ &= 0 + 0 = 0. \end{aligned}$$

The first term of (46) is the difference between the integral of

$$m \left(-I_n f_\nu(\Delta u(\theta, x_a, x_b)) + 4(\theta - x_a)(\theta - x_a)^\top f'(\Delta u(\theta, x_a, x_b)) \right)$$

with respect to ρ^n and with respect to ρ as $n \rightarrow \infty$. The limit is zero componentwise because the integrand is bounded (as we argued above), the integrand is continuous in θ , and ρ^n converges to ρ in the weak topology.

The limit of the second term of (46) is zero because the integrand is Lipschitz continuous in (x_a, x_b) with respect to the distance d , say induced by the L^1 -norm. Lipschitz continuity with Lipschitz constant C implies that we can bound the term by

$$\lim_{n \rightarrow \infty} m \int C d((x_a^n, x_b^n), (x_a, x_b)) d\rho^n = \lim_{n \rightarrow \infty} m C d((x_a^n, x_b^n), (x_a, x_b)) = 0.$$

Lipschitz continuity follows from the gradient of the integrand in (x_a, x_b) being bounded componentwise, which follows analogously to how we showed above that the integrand is bounded.

If $H_a(x_a, x_b, \rho, m)$ is continuous in (x_a, x_b, ρ) , then it is jointly continuous in (x_a, x_b, ρ, m) as m simply multiplies the integrand. Then, $v^\top H_a(x_a, x_b, \rho, m)v$ is jointly continuous in those variables and v .

Using continuity of the Hessian, we show for m small enough, the Hessian is negative definite for $x_a, x_b \in D$ and $\rho \in R$. Recall that a matrix $H \in \mathbb{R}^{n \times n}$ is negative definite if for all $v \in \mathbb{R}^n$, $v^\top H v \leq 0$. Given m , choose the $x_a, x_b \in D$, σ_ρ^2 with $\sigma_\rho^2 \leq 1$, and $v \in \mathbb{R}^n$ with $\langle v, v \rangle = 1$ to maximize $v^\top H_a(x_a, x_b, \rho, m)v$. By the above, $v^\top H_a(x_a, x_b, \rho, m)v$ is continuous in $(x_a, x_b, \sigma_\rho^2, v, m)$ and the choice set is compact. At $m = 0$, we have $H_a(x_a, x_b, \rho, m) = -I_n$, so the value is -1 irrespective of the choice of (x_a, x_b, ρ, v) , so the maximum is also -1 . By Berge's maximum theorem, the value function is continuous, so for some $\underline{m}$ the value function crosses zero for the last time and is negative for $m < \underline{m}$. This implies that for $m < \underline{m}$, $v^\top H_a(x_a, x_b, \rho, m)v < 0$ for all $v \in \mathbb{R}^n$, so the Hessian is negative definite for all $x_a, x_b \in D$ and $\rho \in R$. This implies concavity of $U_a(x_a, x_b, \rho)$ in x_a for all $x_a, x_b \in D$ and all $\rho \in R$.

A similar argument shows that scaling up the valence shock ν by a large enough factor k makes the party objective concave. Valence shock $k\nu$ has the density $\frac{1}{k} f_\nu(\frac{x}{k})$. Reparametrizing by $c = 1/k$, we get the density $c f_\nu(cx)$ and the derivative of the density being $c^2 f'_\nu(cx)$. At $c = 0$, the Hessian is thus $-I_n$, which is negative definite, for all x_a, x_b, ρ . To show for c small enough (or, equivalently, k large enough), the Hessian is negative definite, again, let a fictitious adversarial agent choose $x_a, x_b \in D$, $\rho \in R$, and $v \in \mathbb{R}^n$ with $\langle v, v \rangle = 1$ to maximize $v^\top H_a(x_a, x_b, \rho)v$. Again, by continuity we can apply Berge's maximum theorem to obtain that the Hessian is negative definite for all $x_a, x_b \in D$ and $\rho \in R$ for c small enough.⁴⁰ $\square$

⁴⁰The arguments above use that for m small enough or ν large enough, the party objective becomes dominated by the ideological motive, which is concave, while the vote share motive becomes arbitrarily small. Additionally, one could show for large enough valence shock ν , the vote share alone becomes concave. This follows from the fact that in the preceding paragraph, the density scales by c while the derivative of the density scales by c^2 . Thus, the negative definite density term in the integrand dominates the integrand for small enough c .