

NBER WORKING PAPER SERIES

AMBIGUOUS ATTRIBUTION:
THEORY AND EVIDENCE

Ricardo Alonso
Monica Martinez-Bravo
Gerard Padró I Miquel
Carlos Sanz
Silvia Vannutelli

Working Paper 35550
<http://www.nber.org/papers/w35550>

NATIONAL BUREAU OF ECONOMIC RESEARCH
1050 Massachusetts Avenue
Cambridge, MA 02138
July 2026

We would like to thank Adriano Pace for outstanding research assistance. We also thank Charles Angelucci, Michael Becher, Daniel A. N. Goldstein, Till Nikolka, Maria Petrova and Massimo Pulejo for excellent discussions, and seminar participants at the China and the Global Economy seminar, the NBER Organizational Economics Working Group, LACEA-RIDGE, IESE, Stanford University, UC Berkeley, NYU, and EPCS, EPSS, LE-ifo-CIE, MESS, CEPR Rome Media, and PEDD conferences. The views expressed in this paper are those of the authors and do not necessarily reflect the position of the Bank of Spain, the Eurosystem, or the National Bureau of Economic Research.

NBER working papers are circulated for discussion and comment purposes. They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2026 by Ricardo Alonso, Monica Martinez-Bravo, Gerard Padró I Miquel, Carlos Sanz, and Silvia Vannutelli. All rights reserved. Short sections of text, not to exceed two paragraphs, may be quoted without explicit permission provided that full credit, including © notice, is given to the source.

Ambiguous Attribution: Theory and Evidence

Ricardo Alonso, Monica Martinez-Bravo, Gerard Padró I Miquel, Carlos Sanz, and Silvia Vannutelli

NBER Working Paper No. 35550

July 2026

JEL No. D83, H19, H49, H7, P16

ABSTRACT

Clarity of responsibility is an essential element of political accountability. We develop a rational model of Bayesian updating in the presence of ambiguous attribution and we test its predictions using an original survey. We show that respondents' partisanship, assessment of public healthcare quality, and beliefs over which layer of government is responsible for healthcare are correlated as predicted: good-assessment voters attribute responsibility to the layer governed by their preferred party, while bad-assessment voters blame the layer governed by the party they dislike. These partisan patterns of credit and blame, often interpreted as evidence of motivated reasoning or partisan bias, can thus arise from rational Bayesian updating under attribution ambiguity. No such partisan patterns exist where the same party is in charge of regional and central government. A survey experiment in which we inform subjects of the official quality of healthcare has them update in the predicted, partisan, direction. Model and empirical results show that partisan priors are extremely hard to dislodge when attribution is ambiguous.

Ricardo Alonso
London School of Economics
and Political Science (LSE)
Department of Management
R.Alonso@lse.ac.uk

Monica Martinez-Bravo
Centro de Estudios Monetarios
y Financieros (CEMFI)
mmb@cemfi.es

Gerard Padró I Miquel
Yale University
and NBER
gerard.padroimiquel@yale.edu

Carlos Sanz
Banco de España
carlossanz@bde.es

Silvia Vannutelli
Northwestern University
and NBER
silvia.vannutelli@northwestern.edu

1 Introduction

Democratic accountability presupposes that voters observe outcomes, correctly attribute responsibility for what they observe, and update their beliefs and voting behavior accordingly (Fearon, 1999; Ashworth, 2012a; Malhotra and Kuo, 2008). In practice, it is often difficult to assign responsibility. For most policy outcomes — a slow public hospital, a commuter-train breakdown, rising inflation, a failing school, a botched disaster response — it is typically unclear which political actor took the decisions that produced the outcome. For example, current inflation or budget deficits can be blamed on decisions taken by administrations long past. Certainly, no individual voter has direct access to the underlying policy production function. Therefore, ambiguity of attribution is not an edge case but a pervasive feature of modern policy-making, and it sits upstream of everything a rational voter does with information on the quality of government.

We focus on one particularly consequential source of attribution ambiguity: the allocation of responsibility across layers of government. Multilayer governance is ubiquitous: in federal systems such as the United States, Germany, or Spain; in decentralized but non-federal states, where subnational governments administer centrally funded services; and in supranational entities such as the European Union, where policy failures may genuinely be the fault of Brussels or of national governments. In each of these settings, areas of responsibility overlap in ways that are difficult for voters to parse. How citizens form and revise beliefs about who is responsible in such environments — and how their partisan sympathies interact with this inferential problem — is the central question of this paper.

We start by developing a simple rational-Bayesian model of belief updating under attribution ambiguity. Two political actors (henceforth parties, which we refer to as Right and Left) are plausibly responsible for a policy outcome. Citizens are uncertain about three objects: the competence of Right, the competence of Left, and which party is responsible for the policy outcome. Each citizen’s experience with the policy is privately observed, and the probability of obtaining a good outcome is increasing in the competence of the party which is, in fact, responsible. Upon observing the outcome, the citizen must simultaneously update over the three objects she is uncertain about. In this simple framework, partisanship has a minimalistic conceptualization: a citizen is a left-wing (right-wing) partisan if, according to her priors, the Left (Right) party is

more competent on average.

Our first result indicates that updating over attribution is a *necessarily* partisan exercise. Observing the same outcome, citizens of differing partisanship rationally update in opposite directions about who is likely responsible: after a good outcome, each partisan credits her preferred actor; after a bad outcome, each shifts blame to the party they do not support. The intuition is very stark: say a right-wing citizen observes a success. She is uncertain about which party is responsible, but, by virtue of her priors, believes that the Right party is more likely to produce successes. She thus cannot help but update upwards her belief that the Right party is responsible. In contrast, a left-wing citizen who also observes a success will update downwards her belief that the Right party is responsible. Therefore, this simple rational model generates patterns that closely resemble what the literature has interpreted as behavioral partisan distortions.¹

The second result shows that updating over party competence is a difficult exercise in the presence of ambiguous attribution. The reason behind this complication is very intuitive: a citizen who observes a success but is unsure who has produced it, necessarily must update upwards on the competence of *both* parties. Given this, whether the partisan gap (her belief over the difference between the competence of Right and Left) increases or decreases is inherently ambiguous. More specifically, we show that the direction of updating over the gap depends on the prior over attribution and on the citizen's relative uncertainty over the quality of both parties, a conceptually clear but empirically slippery object. It follows from this result that information on performance in the presence of attribution ambiguity can easily have muted or counterintuitive effects on citizens' partisanship. This rationalizes the somewhat disappointing empirical findings of a large experimental literature which investigates the effect of information treatments on political outcomes.²

The third result elucidates the empirical implications of this theory to identify conditions under which beliefs over attribution, partisanship, and experience of policy success are correlated

¹The fact that citizens appear eager to credit their preferred side for good outcomes and blame the opposing side for bad outcomes is often interpreted as evidence of partisan bias, motivated reasoning, or selective attribution: (Healy et al., 2014; Bisgaard, 2019; Graham and Singh, 2024). Our point is that this pattern is not, by itself, diagnostic of behavioral bias.

²A well-known meta study concludes that “we find no evidence overall that typical, nonpartisan voter information campaigns shape voter behavior” (Dunning et al., 2019).

at the aggregate level. More specifically, we show that as long as partisanship is minimally stable, we should see the following structure of beliefs: citizens who assess the quality of policy positively and are right-wing (left-wing), should attribute the policy to Right at a higher (lower) rate compared to citizens who are bad-assessment and right-wing (left-wing). Similarly, among good-assessment (bad-assessment) citizens, those who are right-wing should attribute the policy to Right at a higher (lower) rate compared to citizens who are left-wing.

We take these predictions to the data studying beliefs over the Spanish public healthcare system, which offers an unusually clean setting for studying attribution ambiguity. Spain is a highly decentralized state in which regional governments (the Autonomous Communities) are constitutionally responsible for organizing and delivering healthcare, while the central government sets the overall legal framework, coordinates inter-regional policy, and provides much of the financing. Both layers can plausibly be responsible for observed performance, and in fact both are routinely blamed in public discourse.

Our design follows a growing survey-based literature that combines detailed belief elicitation with randomized information provision to study how individuals reason about complex policy environments (Stantcheva, 2021; Haaland et al., 2023). Specifically, between July 10 and 17, 2023, in the week immediately before the Spanish general election, we fielded an original survey on a representative sample of 3,857 adults. The survey elicits, for each respondent: partisan identity; beliefs about which layer of government is responsible for healthcare, separately for legal (*de jure*) and political (*de facto*) responsibility; and beliefs over how long it takes on average in their region of residence to see a public health medical specialist and to receive elective surgery. The latter variables are our proxy for a respondent’s assessment of the quality of public healthcare: shorter waiting lines are a vertical good which should be the result of higher quality management and better husbandry of public resources.

We first document that there is wide dispersion of assessments of quality. Despite the fact that average waiting times are published online by the regional healthcare ministries, for each region and partisanship, we find substantial overestimation and underestimation among respondents: more than one-fifth of respondents hold beliefs more than a standard deviation from the regional official waiting time. Second, respondents are informed about the legal allocation of healthcare responsibility: approximately 60% correctly identify the regional level as legally re-

sponsible for healthcare management. Political attribution, however, is far more dispersed — 22% assign primary responsibility to the central government, 36% to regional governments, with the remainder distributed in between. The gap between accurate beliefs about legal competence and dispersed beliefs about political attribution is telling: attribution is not hard because voters lack institutional knowledge, but because political responsibility is genuinely ambiguous.

We then proceed to map out the structure of respondents' beliefs. We show that despite the huge aggregate heterogeneity in assessment of quality and attribution, beliefs are correlated as predicted by the model. At the time in Spain the central government was controlled by Left. When we look at right-wing respondents who live in regions governed by Right, we find that those who believe waiting times are short (i.e. healthcare is high quality), are more likely to assign responsibility to the regional government than those who believe healthcare quality is low. The opposite is true for left-wing voters. Similarly, when we look at voters who believe healthcare is high quality, those who are right-wing are more likely to assign responsibility to the regional government than those who are left-wing. The opposite is true among respondents who believe healthcare waiting lines are long.

Importantly, we can check whether these patterns exist in regions where the regional government is Left. Note that in these regions the partisan cue is not present, as both the central and the regional government are run by the same party. Indeed, we find that in these regions the aforementioned patterns are either muted or entirely absent. In other words, while institutional ambiguity remains, respondents cannot use party quality as an updating factor and thus the correlation between partisanship, beliefs over healthcare quality, and attribution is not present. We can thus use these regions both as placebo and also as control in triple-difference specifications to purge the spurious effects of unmodeled reasons why ideology or perceptions of quality may affect attribution beliefs. The theory-predicted patterns in regions governed by Right emerge in full force in this last set of specifications.

According to the model, the correlational structure of respondents' beliefs is to be interpreted as the effect of partisanship and assessment of quality on attribution of responsibility. However, an existing literature (e.g., Bullock et al., 2015) documents an alternative causal channel: a partisan respondent who believes her preferred party is responsible for healthcare may report she believes healthcare is high quality – an expressive phenomenon dubbed partisan cheer-leading, which

would amount to reverse causality in our setting. To ensure our preferred interpretation is a major force behind the correlational structure we have described, we embedded an experiment in the survey: two-thirds of respondents were randomly assigned to receive accurate, region-specific information on the official waiting times, while one-third served as a pure control. We thus shock voters' beliefs about performance exogenously — using a piece of information that is hard to dismiss as politically biased because it is numeric, region-specific, and comes from official statistics — and observe how attribution of responsibility responds.

The informational treatment induces updating in the expected, heterogeneous, direction. Given identical numeric news about official average regional waiting times, left- and right-wing voters revise attribution in opposite ways, and the direction of the update flips between positive and negative news consistent with the model's sign restrictions. Importantly, nonparametric learning curves show that, across all partisan-alignment groups, posterior beliefs about the performance variable itself move toward the official waiting time: respondents do update in a Bayesian manner. In other words, we can reject the specific irrational bias of stubborn, prior-preserving dismissal of dissonant information. Instead, we find support for our theoretical model which predicts that attribution beliefs do not converge even if extremely precise information about performance is provided.

Taken together, these results show that patterns often interpreted as partisan distortions or motivated reasoning arise naturally when voters are fully rational but face non-trivial uncertainty about who is responsible for outcomes. This has immediate consequences for how we think about political and institutional accountability. When responsibility is ambiguous, the sign of any citizen's response to a given performance signal — whether the decision is a vote, a complaint to her representative, or a consequential fiscal choice such as paying taxes — depends on the interaction of her partisan priors and her attribution beliefs. The overarching lesson here is that even if voters were fully rational, in the presence of attribution uncertainty, information about government performance, no matter how abundant, is unlikely to dislodge their partisan priors.

Our framework and evidence thus provide a new lens to revisit a large empirical literature in political economy, development economics, and public finance. In recent decades, many experiments have informed citizens about government performance and measured effects on a wide range of outcomes including voting, complaint-making, and tax compliance. The results are strik-

ingly heterogeneous: large and positive in some settings, null in many, and occasionally of the opposite sign (Ferraz and Finan, 2008; Dunning et al., 2019; Incerti, 2020; Chong et al., 2015; Banerjee et al., 2024; Bhandari et al., 2023; Adida et al., 2017; Grossman and Michelitch, 2018; Brockmeyer et al., 2024; Hurst et al., 2024). This can be rationalized as a direct prediction of rational inference under attribution ambiguity, with no appeal to voter irrationality required: in a world of genuinely ambiguous attribution, the effect of performance information on electoral advantage is not theoretically pinned down, so null and mixed findings are predictions of the rational model rather than anomalies to be explained. A methodological implication follows: experiments that inform citizens about performance without measuring and conditioning on attribution beliefs are estimating a reduced-form parameter whose sign depends on an unobserved input — the partisan structure of voters’ attribution beliefs — that this literature has rarely measured.

We also contribute to the theoretical literature on political agency and electoral accountability (Ashworth, 2012b; Ashworth et al., 2017; Besley and Prat, 2006), which typically assumes that voters observe outcomes and attribute them to specific incumbents even if there is still an inference problem caused by hidden information. In ambiguous settings this assumption is often unrealistic, and our framework studies the rational consequences of such ambiguity. More broadly, our results speak to the active debate in political economy about whether apparent partisan distortions reflect motivated reasoning or differing priors.³ Relative to the broader tradition of eliciting partisan-structured perceptions of macroeconomic and policy quantities (Bartels, 2002; Bisbee and Lee, 2022; Bisgaard, 2019; Bonomi et al., 2021), we measure not only outcome perceptions but also the attribution beliefs they map into, and we show that partisanship acts primarily through

³Little et al. (2022) and Little (2025) probe the observational equivalence between once-motivated-reasoning models and Bayesian models with different priors. Gentzkow et al. (2025) similarly shows that fully rational, Bayesian voters can depart from standard updating in the presence of uncertainty about the ideological bias and trustworthiness of information sources — a different source of uncertainty than the one we study, but parallel in its effects: rational agents who lack a key piece of information necessary for clean inference produce belief patterns that look at first glance like irrational distortion. These results are the theoretical counterpart of our empirical finding that rationality plus uncertainty about responsibility reproduces the patterns typically read as motivated reasoning. Independent recent evidence supports this interpretation: Graham and Singh (2024) document that COVID-19 selective attribution in the United States is driven primarily by updating of competence beliefs rather than by motivated reasoning.

the attribution channel rather than through distortion of outcome beliefs themselves.

The observation that multilayer governance blurs the link between outcomes and political actors has a long descriptive tradition in political science (Powell and Whitten, 1993; Anderson, 2006; León, 2011; León et al., 2018; Hobolt and Tilley, 2014; Heinkelmann-Wild et al., 2023). A growing experimental literature complements this tradition with survey and field experiments designed to move attribution beliefs directly — through partisan labels, responsibility cues, corrections about legal competence, or variation in the perceived balance of power between politicians and bureaucrats (Tilley and Hobolt, 2011; Malhotra and Kuo, 2008; León and Orriols, 2019; Fernández-Albertos et al., 2013; Martin and Raffler, 2021; Capezzone et al., 2026). A common feature of these designs is that the researcher manipulates beliefs about *who is responsible* rather than beliefs about the underlying performance outcome. This has delivered suggestive descriptive evidence that attribution beliefs can be moved and that they mediate evaluations of public services. For our questions, however, it has two limitations. First, in multilayer settings political responsibility is genuinely contested — no researcher can credibly reveal the “true” allocation of central versus regional authority. Second, because pre-treatment beliefs on the shocked variable are typically not elicited, within-person belief revision cannot be observed, and sharp rational-updating hypotheses cannot be tested directly. Instead, we shock perceptions of quality using numeric information drawn from official statistics which is harder for subjects to dismiss as politically motivated statements about who is responsible. Moreover, we elicit prior and posterior beliefs on both performance and attribution, which lets us trace within-person updating non-parametrically, test the model’s full sign structure on heterogeneous treatment effects, and rule out the specific irrationality — stubborn, prior-preserving dismissal of dissonant information — that is sometimes invoked to explain persistent partisan beliefs. We do not claim novelty on the underlying empirical regularity — that attribution in multilayer settings is dispersed and partisan-structured — but on the framework within which it is interpreted and on the design that tests it.

2 Theory

In this Section we introduce a stylized model of citizen learning to explain our empirical results and explore their implications. In our model, citizens privately experience a public service which

is implemented by one of two politicians. Whether the experience is a success depends on the competence/ability of the responsible politician, but citizens do not know which politician is to be held responsible. Hence, they use the observed experience to both attribute responsibility and learn about politicians' abilities.

Politicians There are two politicians (political parties) L and R that are currently in government and hence are plausibly in charge of a public service. Politicians differ in their skills: Politician $j \in \{L, R\}$ has skill $\theta_j \in [0, 1]$. We let $a \in \{0, 1\}$ denote which politician is truly responsible for the public service, with R (L) implementing it if $a = 1$ ($a = 0$).

Citizens There is a countably infinite mass of citizens, indexed by $\mathcal{I}$. They are unsure about both a) the skill of politicians and b) who is responsible for the public service. Citizens may have different prior beliefs over $[a, \theta_R, \theta_L]$, and we assume such priors are independent across all three random variables, i.e., for $i \in \mathcal{I}$

$$F[a, \theta_R, \theta_L; i] = F_a^i(a)F_R^i(\theta_R)F_L^i(\theta_L).$$

We let $\alpha_i(a) \equiv \Pr_i[a]$ and we assume that $F_j^i, j \in \{R, L\}$, admits a continuous density $f_j^i > 0$.

Private Experience Each citizen has a private experience of the public service. We model this as a service outcome $Y_i \in \{F, S\}, i \in \mathcal{I}$, which can be a success S or failure F . Outcomes are conditionally independent across the population but depend on the skill of the responsible politician. More specifically, the conditional probability of success for each citizen is

$$\Pr[Y_i = S | a, \theta_R, \theta_L] = a\theta_R + (1 - a)\theta_L.$$

Therefore, while there is a single politician responsible for the policy, there is variation in the experience perceived by citizens. We thus have two reasons for heterogeneity in citizens' posteriors over responsibility: heterogeneity in priors and variation in experience.

2.1 Properties of Citizens' Inference

It is useful to divide citizens into two groups, $\mathcal{R}$ and $\mathcal{L}$, according to their priors. Voters in group $\mathcal{R}$ initially believe that R is more competent on average than L —so, $\mathcal{R} \equiv \{i \in \mathcal{I} : \mathbb{E}_i[\theta_R] \geq$

$\mathbb{E}_i[\theta_L]$ — and we will say that R is their favorite party. For voters in $\mathcal{L}$, L is their favorite party— so, $\mathcal{L} \equiv \{i \in \mathcal{I} : \mathbb{E}_i[\theta_R] < \mathbb{E}_i[\theta_L]\}$. These definitions capture partisanship in this simple model: that a voter is partisan simply means that she believes that her favorite party on average produces more competent politicians than the opposite party.

The first result we present concerns how rational voters update over attribution as a function of prior partisanship and experience with the public service.

Proposition 1. *Suppose that citizen i experiences outcome $Y_i = Y$ and let $p_i(Y) \equiv \Pr_i[a = 1 | Y_i = Y]$ be this citizen's posterior over a . Then:*

- $p_i(Y = S) \geq \alpha_i(a = 1)$ if $i \in \mathcal{R}$.
- $p_i(Y = S) \leq \alpha_i(a = 1)$ if $i \in \mathcal{L}$.
- $p_i(Y = F) \leq \alpha_i(a = 1)$ if $i \in \mathcal{R}$.
- $p_i(Y = F) \geq \alpha_i(a = 1)$ if $i \in \mathcal{L}$.

This result is very stark. If a citizen experiences a success, she necessarily concludes that her favorite party is more likely to be responsible for the service than she previously thought. The opposite happens if she experiences a failure. This is intuitive: since she believes one party to be better at generating successes than the other, she must take a success as evidence that her favorite party is more likely to be responsible. Notice that a direct implication of this result is that, sharing the same experience, two citizens with different partisanship update in opposite ways regarding who is responsible for the service. In summary, voters attribute successful outcomes to their favorite party, updating in opposite directions based on their partisan priors. And this is fully rational.

While the updating process over a is clear as a function of partisanship and outcome, ambiguity over responsibility causes updating over party quality to be more nuanced. To see this, let $\mathbb{E}_i[\theta_R - \theta_L]$ be the advantage that party R has over L according to citizen i 's prior. Upon observing Y_i , the voter updates on both θ_R and θ_L . This results in an updated advantage for party R over L , $\Delta_R(Y_i) \equiv \mathbb{E}_i[\theta_R - \theta_L | Y_i]$. Our second result characterizes how the advantage evolves as voters update.

Proposition 2. Let $\tilde{\Delta}_R^i(Y_i) = \Delta_R(Y_i) - \mathbb{E}_i[\theta_R - \theta_L]$ be the change in advantage as voter i observes Y_i . We have:

- Suppose all voters know $a = 1$. Then $\tilde{\Delta}_R^i(Y_i = S) > 0$ and $\tilde{\Delta}_R^i(Y_i = F) < 0$ for all i .
- Suppose all voters know $a = 0$. Then $\tilde{\Delta}_R^i(Y_i = S) < 0$ and $\tilde{\Delta}_R^i(Y_i = F) > 0$ for all i .
- When $\alpha_i(a) \in (0, 1)$, $\tilde{\Delta}_R^i(Y_i = S) > 0$ and $\tilde{\Delta}_R^i(Y_i = F) < 0$ iff

$$\frac{\alpha_i(a = 1)}{\alpha_i(a = 0)} > \frac{\text{Var}_i[\theta_L]}{\text{Var}_i[\theta_R]}. \quad (1)$$

This proposition shows there is a clear difference in the scenario with certainty over responsibility versus the scenario with ambiguity. The first two results concern the unambiguous scenario and are simple and intuitive: if citizens know that R is responsible for the policy, then Y_i contains no information about θ_L . Accordingly, observing S must lead to a positive update on θ_R and an increase in R 's advantage *independently* of citizen i 's partisan priors. This scenario thus bodes well for accountability: the advantage of the party known to be responsible for the service increases with successes and shrinks with failures.

The last result in the proposition shows that when there is uncertainty over responsibility, updating on party advantage depends on the uncertainty faced by the citizen on *both* politicians as captured by the second moment of her priors over party quality. To understand this, notice that upon observing S (F), a citizen must update upwards (downwards) on the quality of *both* parties. This is indeed a direct consequence of not knowing for certain who implements the service. Therefore, what happens to the advantage of party R depends on which θ_j updates the most. Naturally, if voter i thinks that R is more likely to be responsible than L , she will tend to update θ_R more. This is captured by the left hand side of (1). However, if she is a lot more uncertain about the quality of L than she is about the quality of R , her beliefs over θ_L are a lot more reactive to new information than her beliefs over θ_R . This is the reason why the relative uncertainty of the priors on quality matters, as captured in the right hand side of (1).

There are two important implications of this result. First, when voters are uncertain about attribution, accountability is necessarily muddled by factors extraneous to policy performance. In particular, an incumbent could be doing a great job, providing S to most of the population, and

still see his electoral advantage shrink because most voters happen to be very uncertain about the quality of his challenger.

Second, this model does not necessarily predict rigid partisanship. Proposition 1 suggests that the advantaged party should benefit from ambiguity as successes are interpreted as being likely produced by them. However, this could happen at the same time that the advantage shrinks if the second moment of a voter's beliefs does not satisfy (1).

To see more concretely how ambiguity over responsibility can work against accountability, consider a simple numerical example. Suppose voter $i \in \mathcal{R}$ initially assigns average abilities $\mathbb{E}_i[\theta_R] = 0.6 > 1/2 = \mathbb{E}_i[\theta_L]$ and believes that party R , her favorite, is somewhat more likely than L to be responsible for the service, with $\alpha_i(a = 1) = 0.7$ and $\alpha_i(a = 0) = 0.3$, so that

$$\frac{\alpha_i(a = 1)}{\alpha_i(a = 0)} = \frac{0.7}{0.3} \approx 2.33.$$

At the same time, she is quite confident about the quality of R but very uncertain about the quality of L . For instance, her prior variance over θ_R is relatively small, say $\text{Var}_i[\theta_R] = 0.02$, while her prior variance over θ_L is maximal given the average, i.e., $\text{Var}_i[\theta_L] = 0.5 * (1 - 0.5) = 0.25$, so that

$$\frac{\text{Var}_i[\theta_L]}{\text{Var}_i[\theta_R]} = \frac{0.25}{0.02} = 12.5.$$

In this case, condition (1) fails, because

$$\frac{\alpha_i(a = 1)}{\alpha_i(a = 0)} = 2.33 < 12.5 = \frac{\text{Var}_i[\theta_L]}{\text{Var}_i[\theta_R]}.$$

Now suppose voter i experiences a success, $Y_i = S$. By Proposition 1, she rationally concludes that R is more likely to be responsible than she previously thought, so $p_i(Y_i = S) > \alpha_i(a = 1)$. At the same time, however, Proposition 2 implies that her perceived advantage of R over L *shrinks*, $\tilde{\Delta}_R^i(Y_i = S) < 0$. Intuitively, because she is much more uncertain about θ_L than about θ_R , the same success leads her to revise her beliefs about L 's competence more than about R 's competence. In fact, her uncertainty over L 's competence was so large, that R 's posterior advantage after observing $Y_i = S$ turns negative, as $\Delta_R(Y_i = S) = -0.007$, and she becomes an L -supporter.⁴ In this case, the success makes it more likely in her mind that R delivered

⁴We formally show this in the proof of Proposition 2. This is an example of a *crossover* voter which we define below.

the service, but it also narrows the perceived quality gap to an extent that she switches party allegiance.

This example illustrates the tension highlighted above. With ambiguous responsibility, the same good performance can simultaneously (i) strengthen the posterior belief that the incumbent is responsible and (ii) weaken the perceived quality advantage of the incumbent over the challenger. In such environments, even a government that delivers many successes may fail to build, or may even lose, a quality advantage in voters' eyes if voters' prior uncertainty about the challenger is sufficiently large.

We now turn from individual-level updating to the empirical content of the model. In particular, we derive predictions for how, in the population, beliefs about attribution should co-vary with partisanship and with citizens' private assessments of service quality. These predictions motivate both the structure of our survey and the hypotheses we test in the control and treatment groups.

2.2 Empirical Implications

The model above induces a specific structure on the *population* correlation between partisanship, private outcomes, and beliefs over attribution. The following proposition provides bounds on this correlation in terms of the largest group of voters that may revise partisanship in light of their private outcome. To present this result, let $A_j(Y) \in \{0, 1\}$, $j \in \{R, L\}$, be equal to 1 if a citizen believes after observing outcome Y that politician j is more likely to be responsible for the policy.⁵ Furthermore, partition the initial supporters of R and L into the following two groups according to politicians' expected skill. First, the base of party j , $base_j(Y)$, comprises those citizens who are strong supporters of party j :

$$base_j(Y = S) \equiv \{i \in \mathcal{I} : (\mathbb{E}_i[\theta_j])^2 \geq \mathbb{E}_i[\theta_k]\},$$

$$base_j(Y = F) \equiv \{i \in \mathcal{I} : (\mathbb{E}_i[1 - \theta_k])^2 \geq \mathbb{E}_i[1 - \theta_j]\}.$$

This definition only considers the first moment in the distribution of priors, which leaves substantial degrees of freedom for citizens' learning to vary with their prior $F^i(\theta_R, \theta_L) = F_R^i(\theta_R)F_L^i(\theta_L)$.

⁵Formally, we let $A_R(Y) = \mathbb{I}\{p(Y) \geq 1/2\}$ and $A_L(Y) = \mathbb{I}\{p(Y) \leq 1/2\}$.

Nevertheless, we show in the proof of Proposition 3 that any citizen in $base_j(Y)$ will remain a supporter of party j after observing outcome Y *irrespective* of the variance they face on the skill of the politicians.

Potential crossover supporters of party j , $crossover_j(Y)$, are those that may switch allegiance after observing Y . The formal characterization of crossover citizens is therefore just a combination of sets: citizens who are supporters today, but not in the base.⁶ These preliminaries allow for the statement of the main population-level result:⁷

Proposition 3. *Let*

$$\rho_{jk}(Y) = \Pr[A_j(Y) = 1 | \Delta_k(Y) \geq 0, Y]$$

be the fraction of citizens that, after observing Y , support politician k and believe politician j is responsible. Suppose voters have full uncertainty over attribution—i.e., $\alpha(1) = \alpha(0) = 1/2$. Then,

(i) *The probability that a citizen attributes a success to their ex-post favorite politician satisfies*

$$\rho_{jj}(Y = S) \geq \frac{\Pr[i \in base_j(Y = S)]}{\Pr[i \in base_j(Y = S)] + \Pr[i \in crossover_k(Y = S)]}. \quad (2)$$

(ii) *The probability that a citizen attributes a failure to their ex-post laggard politician satisfies*

$$\rho_{jk}(Y = F) \geq \frac{\Pr[i \in base_k(Y = F)]}{\Pr[i \in base_k(Y = F)] + \Pr[i \in crossover_j(Y = F)]}, j \neq k. \quad (3)$$

The assumption of full prior uncertainty over attribution implies that initial support determines posterior attribution: for instance, citizens in $\mathcal{R}$ will believe that R is more likely to be responsible after a success—so $A_R(Y = S) = 1$ —while they will assign a failure to L —so $A_R(Y = F) = 0$ —see Proposition 1.⁸ However, posterior partisanship depends on how much citizens learn about each party and our partitioning into $crossover_j(Y)$ and $base_j(Y)$ captures when learning can and cannot trounce initial advantage. To see this through an example, suppose that all citizens in $\mathcal{R}$ are absolutely certain about θ_R —so that $Var(\theta_R) = 0$. If they are also certain about θ_L —so that $Var(\theta_L) = 0$ —then their private experience would not change the sign

⁶Therefore, crossover voters are $crossover_R(Y) = \mathcal{R} \cap (base_R(Y))^c$ and $crossover_L(Y) = \mathcal{L} \cap (base_L(Y))^c$.

⁷We interpret probabilities over subsets of $\mathcal{J}$ as empirical frequencies over the population.

⁸In the empirical application, we assume we survey citizens after they have had direct or indirect exposure to the public service. Assuming a fully uncertain prior is done for clarity. After all, if one is willing to assume that there was already a correlation in the priors between attribution and partisanship then one can obtain much more easily, and uninterestingly, the correlation in the posteriors.

of R 's advantage. Suppose instead that all citizens in $\mathcal{R}$ are maximally uncertain about θ_L —so that $Var(\theta_L) = \mathbb{E}[\theta_L](1 - \mathbb{E}[\theta_L])$. If $Y = S$, any citizen $i \in \mathcal{R}$ will not revise θ_R but will aggressively increase $\mathbb{E}[\theta_L]$. If $i \in crossover_R(Y = S)$, then R 's initial advantage is modest and the citizen will believe that R is responsible but switch allegiance to L , as in the example in the previous subsection. If instead $i \in base_R(Y = S)$, then the initial *average* advantage is so large that she always keeps her allegiance to R . Thus, (2) puts a lower bound on the fraction of posterior j —partisans that would attribute a success to politician j while (3) provides a lower bound on the fraction attributing a failure to the other politician.⁹

To provide a numerical example, suppose that the cross-sectional distribution of prior beliefs leads to uniformly distributed average abilities of R and L in the population. Then, (2) means that at least $2/3$ of politician- j 's supporters believe him to be responsible after a success, but (3) means that at most $1/3$ believe him to be responsible after a failure. The advantage of Proposition 3 is that it provides this characterization for any distribution of partisan priors in the population.

The following corollary translates the bounds on Proposition 3 into differences in frequencies of attribution across partisanship and experience.

Corollary 1. *Suppose that $\alpha(1) = \alpha(0) = 1/2$ and $\Pr[i \in crossover_j(Y)] \leq \Pr[i \in base_k(Y)]$, for $j, k \in \{R, L\}$. Then,*

(a) *Within partisanship we have*

$$\begin{aligned}\rho_{RR}(Y = S) - \rho_{RR}(Y = F) &\geq 0, \\ \rho_{RL}(Y = S) - \rho_{RL}(Y = F) &\leq 0.\end{aligned}$$

(b) *Within outcomes we have*

$$\begin{aligned}\rho_{RR}(Y = S) - \rho_{RL}(Y = S) &\geq 0, \\ \rho_{RR}(Y = F) - \rho_{RL}(Y = F) &\leq 0.\end{aligned}$$

This corollary details two sets of results we can bring to the survey data we describe below. The first result conditions on partisanship and says that those who report a good assessment of public service (i.e. observe a success in the parlance of the model) are more likely to attribute it to

⁹These bounds will bind when *all* crossover voters switch allegiance as these would minimize the fraction of posterior supporters of a politician that were also initial supporters.

their reported preferred party. In particular, among right-wing voters, those who report a good assessment should assign responsibility to Right at a higher rate than those who do not. In contrast, among left-wing voters, those who report a good assessment should assign responsibility to Right at a lower rate than those who report bad assessment.

The second result conditions on reported assessment instead and says that citizens with a good assessment (bad assessment) attribute responsibility for the policy to their preferred (opposing) party. In particular, good-assessment right-wing voters should assign responsibility to Right at a higher rate than good-assessment left-wing voters. The opposite should be true for those who report bad assessment.

Because observed partisanship is endogenous to experience with the public policy, these associations are not universal due to the issue we report in Proposition 3. The restriction we need to impose on priors is the stated bound on the maximum number of crossover voters relative to both bases. In practice, this assumption requires a very plausible context in which there is reasonable political stability. In other words, the condition is met as long as observing a particular policy outcome does not cause a massive number of citizens to switch party allegiance.

In summary, this simple rational model easily accommodates belief structures in the population in which partisanship, perception of quality of public services, and attribution are correlated, even when individuals only have private signals about service quality. The key point for our purposes is that the model not only rationalizes such patterns, but also constrains their sign and relative magnitude across partisan and outcome groups and thus provides testable implications.

2.3 Does more information resolve attribution ambiguity?

Results so far assume a single round of updating with a private signal. Their implications for learning over attribution are not encouraging as we show that partisanship induces learning in opposite directions. However, one may conjecture that richer information on performance can eventually lead citizens towards the true attribution. The following result establishes that, unfortunately, this is not the case.

Proposition 4. *Let $t = \lim_{n \rightarrow \infty} t(n)$, where $t(n)$ is the cross-sectional policy success rate of the first n citizens. Assume that for every citizen $i \in \mathcal{R}$ ($i \in \mathcal{L}$) the likelihood ratio f_R^i / f_L^i is increasing*

(decreasing). Let $p_i^\infty(t) \equiv \Pr_i[a|t]$ be the (asymptotic) belief over a for voter i . Then:

1. $p_i^\infty(t)$ increases in t if $i \in \mathcal{R}$, but decreases in t if $i \in \mathcal{L}$.
2. For each i , there exist $\tilde{t}_i(Y)$, with $\tilde{t}_i(Y = S) > \tilde{t}_i(Y = F)$, such that they revise their attribution in favor of their preferred candidate after learning t when their private outcome was Y iff $t \geq \tilde{t}_i(Y)$.

This proposition contains two results. The first result demonstrates that priors affect the direction of learning even when information over performance is arbitrarily precise and public. Note that t encapsulates the population rate of service success, as the population grows large. This result shows that providing citizens with this rate (in effect, allowing them to learn the true skill of the responsible politician) still results in partisanship determining the direction of updating over attribution: if the policy looks good on aggregate, people with right-wing priors (in $\mathcal{R}$) are more likely to attribute it to R , while people with left-wing priors (in $\mathcal{L}$) attribute it to L . The intuition for this non-convergence is that the econometric model that voters use is not identified (see Acemoglu, Chernozhukov, and Yildiz, 2016): asymptotically, a citizen can rationalize any observed rate t of successes by either assigning $a = 1$ and $\theta_R = t$ or $a = 0$ and $\theta_L = t$. However, she will use her prior belief to infer attribution from service performance in accordance with her partisanship. Therefore, ambiguous attribution muddles inference and prevents agreement.

This result has an empirical implication we will test experimentally: providing exogenous directional information (i.e. a positive or a negative informational shock) on performance should cause updating about attribution to move in opposite directions depending on partisanship.

The second result points to another empirical implication that we test. When we inform citizens of a given partisanship of the average success rate, those who are given a positive shock should update their beliefs on attribution in opposite direction to those who are given a negative shock.

3 Context, Data and Experimental Design

In the rest of the paper we explore the hypotheses set out in the theory in a specific context of multi-layered governance, Spain, and a specific policy domain, public health provision. This

section describes the context as well as the data collection exercise.

3.1 Decentralization in Spain

Contemporaneous Spain is a decentralized country, both politically and in terms of public service delivery. There are seventeen regional governments called *Comunidades Autónomas*. Among the most important public services managed by the regional governments are education and health. Regional governments are responsible for the day-to-day delivery of these services as well as for specific rules, regulations, capital and labor expenditures. However, this autonomy is exercised within the legal framework established by the central government and within funding strictures determined by rules imposed by the center. As a consequence, there is unavoidable ambiguity in assigning responsibility for unsatisfactory outcomes. For example, if schooling or healthcare outcomes are regrettable, is it because of managerial incompetence of the regional government or insufficient funding provided to the region by the central government?

From an electoral point of view, the institutional structure allows citizens to hold the two levels of government independently accountable. This is because there are regional parliamentary elections that are conducted separately from national level parliamentary elections.

This context therefore provides a fruitful setting to test the hypotheses detailed above. Attributing responsibility over public services is a genuinely difficult task as there is ambiguity built into the institutional framework. Electoral variation means that at any point in time there are regions which are ruled by the same coalition as the central government, and regions which are ruled by the opposition. The effect of *partisan priors* on attribution should only be present in the latter, so the former can serve as both placebo and control.¹⁰

Regarding public services, we focus on healthcare provision. An advantage of the setting is that, as opposed to the familiar USA context, an overwhelming majority of citizens in Spain use the public healthcare system and thus should have direct or indirect experience of its quality (Centro de Investigaciones Sociológicas, CIS).¹¹ One salient and objective marker of such quality

¹⁰To be clear, residents of regions in which the regional and central governments are ruled by the same party also face ambiguous attribution of responsibility. However, in this case there should not be a partisan dimension to how they attribute said responsibility between regional and central government.

¹¹Consistent with this, 99% of respondents in our survey report that their health care coverage is publicly funded.

is the length of the waiting lists to undergo an elective surgical procedure or to obtain a visit with a specialist doctor. The number of days it takes is both an objective numerical outcome and an unambiguous marker of quality which is often discussed in the public domain. A respondent's beliefs over this waiting time are therefore an object we can measure and truthfully manipulate.

3.2 Survey

The data used in this project were collected on an online survey that we ran from 10-17 July 2023. Fieldwork was conducted by data analytics firm Yougov. The sampling is designed to be representative of the Spanish adult population according to age, gender, region of residence, and education level. This is achieved through a quota-sampling system.

Our survey proceeds as follows:¹²

1. Background Information. This includes sociodemographic questions, vote in the past regional elections, and ideology on a 1-10 scale, where 1 stands for “extreme left” and 10 stands for “extreme right”.
2. Elicitation of beliefs about waiting times. Respondents are randomized into two equally sized groups, G1 and G2:
 - (a) G1 is asked about the average expected wait, in terms of days, for a consultation with a specialist doctor, for instance, for cardiology, neurology, or dermatology.
 - (b) G2 is asked about the average expected wait, in terms of days, for an elective surgery procedure, for instance, for cataract, inguinal hernia, or hip prosthesis.

Additionally, 23.2% report having a privately purchased health insurance policy (either individually or through a family member), and 9.4% have one provided by their employer. Typically, when families have access to private healthcare, they use it as a supplement, not as a substitute, to the public system.

¹²The design closest to ours is Martinez-Bravo and Sanz (2025), who conduct a Spanish survey experiment in November 2020 on regional COVID-19 contact-tracer capacity: they shock respondents with the true numeric count and measure both prior and posterior beliefs. We extend their approach to a chronically salient public service outside a pandemic context, measure the full three-variable belief structure (partisanship, performance and attribution) and ground the exercise in a formal model whose comparative statics we can directly test.

3. Randomized provision of information about waiting times in the respondent's region of residence.¹³
4. Elicitation of beliefs over responsibility attribution and over the expected wait omitted in step 2:
 - (a) G1 is asked about the average expected wait for an elective surgery procedure.
 - (b) G2 is asked about the average expected wait for a consultation with a specialist doctor.

Appendix E contains the questionnaire.

Our survey design has two advantages for measuring belief updating. First, by eliciting a pre-treatment belief about the waiting-time dimension for which respondents subsequently receive information, we can characterize the size and direction of the information shock at the individual level. This allows us to distinguish respondents who receive good news from those who receive bad news, to test heterogeneous treatment effects by prior beliefs, and to estimate learning rates: the extent to which posterior beliefs move toward the information provided. Second, we elicit post-treatment beliefs about a closely related but distinct healthcare-performance measure. Respondents who are first asked about specialist consultations are later asked about elective surgery, and vice versa. This follows a common recommendation in information-provision experiments (Fuster et al., 2022a; Haaland et al., 2023): asking respondents to report the exact same belief twice may confuse respondents, make the purpose of the design too transparent, or induce consistency concerns. Asking about a related outcome avoids mechanically repeating the same question while still allowing us to test whether information about one dimension of healthcare quality leads respondents to revise beliefs about another dimension. The assumption is that there

¹³We obtain data from the Information System on Waiting Lists in the National Health System (Sistema de información sobre listas de espera en el sistema nacional de salud) published by the Ministry for Health, version updated on April 10, 2023. The waiting time for specialist doctors includes the following: gynecology, ophthalmology, traumatology, dermatology, otorhinolaryngology, neurology, general surgery, urology, gastroenterology, cardiology. The waiting time for surgery procedures includes: cataract, inguinal/femoral hernia, hip prosthesis (or hip replacement), arthroscopy, varicose veins of the lower limbs, cholecystectomy, hallux valgus (bunion), adeno-tonsillectomy, benign prostatic hyperplasia, pilonidal cyst, carpal tunnel syndrome, knee prosthesis (or knee replacement), valvular heart surgery, coronary bypass, hysterectomy.

is an underlying common component of regional healthcare performance which is correlated with both specialist and surgery waiting times.¹⁴

Description of Experiment The informational treatment is as follows:¹⁵

- Participants receive information on the official average waiting times in their region (Autonomous Communities in Spain). The information they receive pertains to the beliefs they were asked in step 2 above. Namely:
 - If respondent is in G1, she receives information about official average waiting time for a consultation with a specialist
 - If respondent is in G2, she receives information about official average waiting time for elective surgery procedures

We assign individuals to the treatment group according to a stratified randomization procedure. First, individuals are classified in 204 groups or strata with similar baseline characteristics in terms of age (3 levels: 18-35 years old, 36-50, 50+), education (2 levels: secondary schooling or less, more than secondary schooling), region (17 levels, one per region), and ideology (2 levels: 1-4 or 5-10 on the ideology scale). Within each stratum, we randomly assign 2/3 of individuals

¹⁴We did not provide accuracy incentives for the waiting-time belief elicitation. This choice was motivated by three considerations. First, our object of interest is respondents' naturally held beliefs about public healthcare performance, rather than their ability to search for official statistics during the survey. Because regional waiting-time statistics are publicly available online, accuracy incentives could have encouraged respondents to look up the correct answer, thereby changing the informational environment we aim to study. Second, incentive schemes for quantitative belief elicitation can be difficult to explain in a short online survey and may themselves generate confusion or additional measurement error. Third, existing evidence suggests that incentives often have limited effects on belief reports in non-political forecasting domains, even though they can matter in settings with strong expressive or partisan motives. For these reasons, and to keep the belief elicitation simple and comparable across respondents, we elicited beliefs without accuracy incentives.

¹⁵The experiment includes two variants of treatment. Specifically, half of the treatment group receives, in addition to the information described here, information about how the relevant waiting time in their region ranks among all regions in Spain. As shown in Section 5.4, we do not find any statistically relevant effect of this additional information and thus, for the purposes of this article, we collapse the two variations into a single treatment group.

to treatment and 1/3 to control. Subgroups G1 and G2, which determine question ordering, are evenly split across both the treatment and control arms.¹⁶

Our initial sample covers 4,620 individuals. Following our pre-analysis plan, we drop individuals who did not pass the attention check (555 observations) or completed the questionnaire in less than 11 minutes (207 additional observations). We also drop one observation with a very large outlier in waiting time beliefs. Our final sample contains 3,857 observations.

3.3 Empirical Implementation of Theoretical Concepts

Our theoretical predictions pertain to the population correlations across three objects: partisanship, assessment of quality of public services, and attribution of responsibility over said public services. We now detail how we map the variables measured in the survey to the theoretical objects of interest.

Partisanship Spain has a complex multiparty landscape both at the regional and the national level. At the same time, these parties’ ideological affinities result in predictable parliamentary coalitions which give support to government formation at both levels. We classify parties in terms of their presence in the national and regional governments as well as their votes in government formation, and we classify voters in terms of their reported votes for parties. More specifically:

- We define as Left parties (L in the model) the parties which are in control of the central government.
- We define a region as a *united region* if one or more of the Left parties is in the ruling coalition of the regional government.
- We define a region as a *divided region* if none of the Left parties is in the ruling coalition of the regional government. Divided regions are Andalusia, Castile and Leon, Catalonia, Madrid, Galicia, and Murcia—see Table C.1 for further details.

¹⁶Our experimental design was pre-specified in a pre-analysis plan (PaP) that we registered with the AEA Randomized Control Trial Registry on 7 July 2023 (AEARCTR-0011740) - see Appendix D for further details. We also obtained approval from the ethics committee at CEMFI for the survey and experimental design (Application Reference Number 3; Approval date: 18 June 2023).

- In divided regions, a voter is right-wing ($\mathbb{E}_i[\theta_R - \theta_L] > 0$) if in the previous regional elections she voted for one of the parties that voted “yes” to the investiture of the regional government. She is left-wing otherwise. Table C.1 shows the list of parties. Alternative classifications based on which parties are part of the regional executive, or which party has the regional presidency, yield very similar results.
- In united regions, a voter is left-wing ($\mathbb{E}_i[\theta_R - \theta_L] < 0$) if she voted for one of the parties that voted “yes” to the investiture of her regional government, and she is right-wing otherwise.
- Following our pre-analysis plan, for individuals with missing past regional vote information, we use ideological placement on a 1-10 scale, where 1 is extreme left and 10 is extreme right. We code a respondent as right-wing if she positions herself at 5 or above.¹⁷

Experience with Public Services We measure assessment Y_i with a respondent’s belief over waiting time. Respondents are asked to input a specific number of days which they expect it takes on average to be attended in their region of residence. We discretize this variable to align it with the theory. More specifically, we code those subjects who predict a longer waiting time than the official average waiting time in their region as having experienced $Y_i = F$. We call them *bad-assessment* respondents. Those who predict shorter waiting times than the official average waiting time in their region we code as $Y_i = S$ and call them *good-assessment* respondents. Classifying respondents relative to the official average waiting time ensures variation in each region and also links these categories directly with the direction of experimental treatment: by definition, bad-assessment respondents receive a positive informational shock in the experimental arm, while good-assessment respondents receive a negative informational shock.¹⁸

¹⁷In Subsection 5.4, we provide evidence that our results are qualitatively robust when we define partisanship based exclusively either on ideological placement or on past vote.

¹⁸All our results are robust to using other criteria to split the sample between good- and bad-assessment respondents. We consider four alternative definitions. First, we classify a respondent as having a good assessment if he or she is below the median belief, regardless of official average waiting time. Second, we split respondents using a median belief which is defined separately for each combination of region and individual’s ideology. Third, we consider a respondent as good assessment if he or she is below the median in the difference between the stated belief and their ideal waiting time, which we also elicit in the survey. Finally, we consider the continuous variable, defined as the standardized difference in Official Waiting Time - Believed Waiting Time. Please see Subsection 5.4

Attribution of Responsibility We measure attribution as the respondents’ answer to “which layer of government do you believe is more responsible for the quality of public healthcare?”. Respondents can choose a number on a scale from -10 to 10 where -10 indicates “all responsibility lies with the central government” and 10 indicates “all responsibility lies with the regional government.”

As indicated in Table C.2 in the Appendix, our sample closely mirrors the Spanish population in terms of demographic characteristics. The main exception is education: as is common in survey samples, our respondents are more highly educated than the population overall. Appendix Table C.3 shows the summary statistics of our main variables.

4 Empirical Strategy

Our empirical goal is to map the theory’s predictions into a set of estimable sign restrictions.

4.1 Population Correlations

The hypotheses generated by Corollary 1 point us to compare mean attribution beliefs across four groups defined by partisanship (left-wing or right-wing) and assessment (good or bad), measured by beliefs over quality of healthcare. These predictions, of course, only apply to respondents in divided regions. A simple fixed effects regression on data from these regions allows us to recover these means:

$$P_i = \alpha_0 + \alpha_1 O_i + \alpha_2 \mathcal{R}_i + \alpha_3 O_i \mathbf{x} \mathcal{R}_i + \phi_g + u_i. \quad (4)$$

In this specification, an individual’s attribution of responsibility over healthcare to the regional government, P_i , is a function of her partisanship, where $\mathcal{R}_i = 1$ if the voter is right-wing, and her belief over the length of waiting lines, where $O_i = 1$ if the voter has a good assessment. In addition, we consider region g effects to make sure we are comparing individuals within a region.

This specification allows us to test the theoretical hypotheses as follows.

$$\begin{aligned}
H_1 : \quad & \alpha_1 < 0 \\
H_2 : \quad & \alpha_1 + \alpha_3 > 0 \\
H_3 : \quad & \alpha_2 < 0 \\
H_4 : \quad & \alpha_2 + \alpha_3 > 0 \\
H_5 : \quad & \alpha_3 > 0
\end{aligned}$$

These restrictions all come from the same underlying pattern in the theory. H_1 , for example, corresponds to the prediction that among left-wing citizens, those who have a good assessment $O_i = 1$ should be less likely to attribute healthcare to the regional, Right government than those who have a bad assessment $O_i = 0$. It is immediate to verify that H_1 and H_2 correspond to the predictions of Corollary 1 that condition on partisanship and compare across assessment, while H_3 and H_4 correspond to the predictions of Corollary 1 that condition on assessment and compare across partisanship.

As these estimates are simple differences across groups, if there are any reasons outside the model why respondents in one of the four groups systematically differ from those in other groups, our estimates will be confounded. One plausible scenario is that $u_i = \zeta_{\mathcal{R}} + \zeta_O + \epsilon_i$ where $\zeta_{\mathcal{R}}$ and ζ_O may be different from 0. For instance, one plausible conjecture is that $\zeta_{\mathcal{R}} < 0$: in the case of Spain, due to historical reasons, right-wing individuals have a stronger normative belief on the role of the central government which may likely spill over to their factual beliefs. Similarly, there could be partisan differences in the rate of survey response which cause a level difference. It is easy to see that this would bias the tests for H_3 and H_4 as they compare across partisan lines. Or if good-assessment individuals systematically differ from bad-assessment individuals, $\zeta_O \neq 0$, the proposed tests of H_1 and H_2 are suspect. This is the reason why we state H_5 : as this test is equivalent to a difference in differences estimator, it absorbs these level biases. Therefore, while H_5 does not have independent theoretical content, obtaining $\alpha_3 < 0$ would falsify the model.

We summarize the predictions of the model in Table 1.

We also present estimates of regression (4) for united regions. Theory predicts no particular partisan pattern in this case, and hence we can interpret these results as an appropriate placebo test of our theory. In addition, note that united regions can provide independent evidence that

Table 1

Model Predictions in Divided Regions: Observational

	Left-wing	Right-wing	
Bad Assessment	Regional (+)	Central (−)	$H_3 < 0$
Good Assessment	Central (−)	Regional (+)	$H_4 > 0$
	$H_1 < 0$	$H_2 > 0$	$H_5 > 0$

Notes: Summary of the sign restrictions for divided regions implied by Corollary 1. A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of divided regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Bad Assessment” and “Good Assessment” indicate whether respondents’ prior beliefs about waiting times exceed or fall below the official regional average, respectively. Cell entries indicate the predicted direction of attribution in divided regions: “Regional” means responsibility is attributed to the regional government, while “Central” means responsibility is attributed to the central government. Signs indicate the predicted direction of the corresponding hypothesis test.

$\zeta_{\mathcal{R}} \neq 0$ or $\zeta_{\mathcal{O}} \neq 0$. Accordingly, we can use united regions as controls for divided regions following a triple difference specification which should purge our estimates of these group-level effects assuming they are constant across regions.¹⁹

¹⁹More specifically, we run the following triple difference specification:

$$\begin{aligned}
P_i = & \beta_0 \\
& + \beta_1 O_i + \beta_2 \mathcal{R}_i + \beta_3 O_i \times \mathcal{R}_i \\
& + \beta_4 D_{ig} + \beta_5 O_i \times D_{ig} + \beta_6 \mathcal{R}_i \times D_{ig} + \beta_7 O_i \times \mathcal{R}_i \times D_{ig} \\
& + \phi_g + u_i,
\end{aligned} \tag{5}$$

where D_{ig} is a dummy that indicates that respondent i resides in a region g which is divided. Note that, since D_{ig} varies only at the region level, its main effect is absorbed by the region fixed effects and is therefore omitted from estimation. An additional advantage of this specification is that group estimate differences would be correctly estimated under the weaker assumption $u_i = \zeta_{\mathcal{LP}} + \zeta_{\mathcal{LO}} + \zeta_{\mathcal{RP}} + \zeta_{\mathcal{RO}} + \epsilon_i$ provided the group-level biases are common across regions.

4.2 Informational Treatment Effects

In the experimental treatment, we reveal official information about performance that affects good assessment and bad assessment individuals in opposite directions. Good assessment individuals learn that real waiting lines are longer than their estimates, while bad assessment individuals learn the opposite. Thus, bad-assessment respondents receive good news, whereas good-assessment respondents receive bad news. There are two sets of theory-driven hypotheses that apply here.

First, Proposition 1 describes the direction of update on attribution as a function of partisanship and whether the news is good or bad. In particular, in divided regions, right-wing partisans who receive good news should update attribution towards the regional government (run by the Right party), while those who receive bad news should update in the opposite direction. The opposite is true for left-wing partisans.

Using only data from divided regions, we can thus obtain the treatment effect in each of the four groups running the following specification:

$$\begin{aligned}
 P_i = & \gamma_0 + \gamma_1 O_i + \gamma_2 \mathcal{R}_i + \gamma_3 O_i \times \mathcal{R}_i \\
 & + \gamma_4 T_i + \gamma_5 O_i \times T_i + \gamma_6 \mathcal{R}_i \times T_i + \gamma_7 O_i \times \mathcal{R}_i \times T_i \\
 & + \phi_g + u_i,
 \end{aligned} \tag{6}$$

where T_i denotes whether subject i received the informational treatment. The coefficients from this specification induce the following tests

$$\begin{aligned}
 HT_1 : & \quad \gamma_4 < 0 \\
 HT_2 : & \quad \gamma_4 + \gamma_5 > 0 \\
 HT_3 : & \quad \gamma_4 + \gamma_6 > 0 \\
 HT_4 : & \quad \gamma_4 + \gamma_5 + \gamma_6 + \gamma_7 < 0.
 \end{aligned} \tag{7}$$

These correspond to the treatment-effect predictions implied by Proposition 1.

We can, in addition, derive predictions from Proposition 4, which pertain to particular pairs of groups updating in opposite directions. For example, upon receiving the same news (good or bad) left-wing and right-wing partisans should update over attribution in opposite directions.

More specifically, the hypotheses derived from Proposition 4 imply the following empirical implications $HE_1 - HE_5$:

$$HE_1 : \quad \gamma_5 > 0$$

$$HE_2 : \quad \gamma_5 + \gamma_7 < 0$$

$$HE_3 : \quad \gamma_6 > 0$$

$$HE_4 : \quad \gamma_6 + \gamma_7 < 0$$

$$HE_5 : \quad \gamma_7 < 0.$$

For example, HE_3 tests whether upon receiving good news, the gap in responsibility assignment to the regional government between right-wing and left-wing citizens increases. These independently derived predictions from Proposition 4 have an added advantage over $HT_1 - HT_4$: they are more empirically resilient to group-level unmodeled idiosyncrasies in updating.²⁰ In identical fashion as in the previous subsection, HE_5 provides an additional falsification test.

As before, we will also show the outcome of applying this specification on united regions data, where our theory does not predict these particular sign restrictions as partisanship should not affect updating. United regions thus constitute a good placebo for our predicted information experimental effects.

Finally, for consistency, we also show the results of the fourth-difference specification where we look at the difference in treatment effects between divided and united regions. This specification requires that group-level idiosyncrasies in updating are common across regions, which is a somewhat stronger assumption than the bias in levels we consider in the previous subsection.²¹

All specifications shown below include, in addition, region fixed effects. In robustness tests, we also consider specifications with the following controls: age group dummies, education dummies, female, an indicator for having kids, living in a rural municipality, and work status (worker,

²⁰We discuss this robustness in Appendix B.

²¹To be precise, we estimate the fourth-difference specification:

Table 2

Model Predictions in Divided Regions: Experimental

	Left	Right	
Good News	Central ($HT1 < 0$)	Regional ($HT3 > 0$)	$HE3 > 0$
Bad News	Regional ($HT2 > 0$)	Central ($HT4 < 0$)	$HE4 < 0$
	$HE1 > 0$	$HE2 < 0$	$HE5 < 0$

Notes: Summary of the sign predictions for divided regions derived from Proposition 1 and Proposition 4. A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of divided regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Good News” and “Bad News” indicate whether the information treatment provides favorable or unfavorable news about waiting times relative to respondents’ prior beliefs. Signs indicate the predicted direction of the corresponding hypothesis test.

student, retired).

$$\begin{aligned}
P_i = & \delta_0 \\
& + \delta_1 O_i + \delta_2 \mathcal{R}_i + \delta_3 O_i \times \mathcal{R}_i \\
& + \delta_4 T_i + \delta_5 O_i \times T_i + \delta_6 \mathcal{R}_i \times T_i + \delta_7 O_i \times \mathcal{R}_i \times T_i \\
& + \delta_8 D_{ig} \\
& + \delta_9 O_i \times D_{ig} + \delta_{10} \mathcal{R}_i \times D_{ig} + \delta_{11} O_i \times \mathcal{R}_i \times D_{ig} \\
& + \delta_{12} T_i \times D_{ig} + \delta_{13} O_i \times T_i \times D_{ig} + \delta_{14} \mathcal{R}_i \times T_i \times D_{ig} + \delta_{15} O_i \times \mathcal{R}_i \times T_i \times D_{ig} \\
& + \phi_g + u_i.
\end{aligned} \tag{8}$$

5 Results

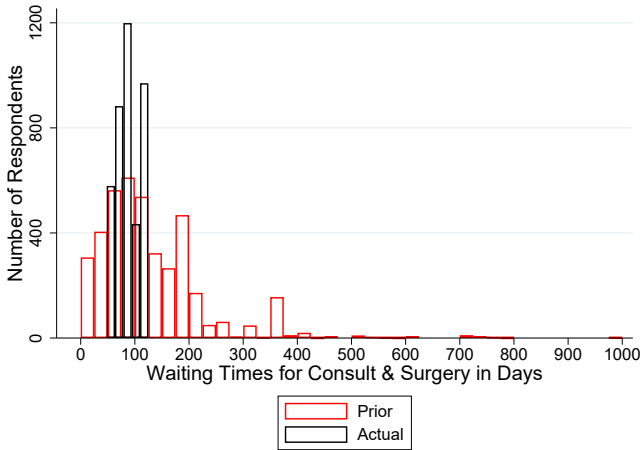
5.1 Do Individuals Have Accurate Information on Waiting Times?

We begin by describing how well citizens understand the performance of the public health system. To do so, we analyze the distribution of subjects' beliefs on waiting times, in Panel A of Figure 1. Beliefs among the control group, depicted in red, range from 0 to approximately 1000 days, with most responses in the interval between 0 and 300 days, and a clear right-skewed distribution with a peak around 50-100 days. In black we show the official average waiting time by region. The official waiting time distribution is much tighter, with a pronounced mode around 100 days, thus revealing substantial misperception.

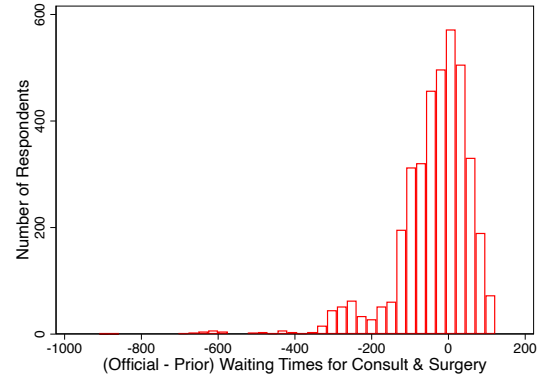
To visualize the magnitude of these misperceptions, for each individual, we compute the difference between the official waiting time in their region and their priors, and we display the distribution of these differences in Panel B of Figure 1. Based on this figure, 60% of respondents can be classified as bad-assessment respondents; these respondents overestimate the time they would need to wait for medical services in their region. This overestimation is visible in the extended right tail of the red distribution in Panel A, where prior beliefs stretch well beyond the official waiting times shown in black. Conversely, 40% of respondents are "good assessment" who underestimate waiting times. This is evident in the left portion of Panel A, where some respondents anticipated shorter waiting times than the official average.

Figure 1: Distribution of Waiting Times

Panel A: Beliefs vs Official Waiting Times



Panel B: Distribution of (Official - Belief)



Notes: Panel A displays the distribution of respondents' prior beliefs about waiting times and official regional waiting times for specialist consultations and elective surgery procedures in days. Panel B displays the distribution of the difference between official waiting times and respondents' prior beliefs. Respondents are classified as "bad assessment" if they overestimate waiting times relative to the official regional average and as "good assessment" otherwise. Sample restricted to control group respondents ($N = 1,274$). Official waiting times are from the Information System on Waiting Lists in the National Health System (Sistema de Información sobre Listas de Espera en el Sistema Nacional de Salud), Ministry of Health, April 2023.

The magnitude of misperception is substantial: approximately 20% of respondents hold beliefs that deviate from the official figure by more than one standard deviation, indicating significant variation in beliefs over the healthcare system's efficiency. The presence of a long right tail also indicates that a subset of respondents holds particularly bad-assessment views about waiting times, expecting delays far longer than the official average.

These findings suggest that public perceptions of healthcare waiting times are characterized by substantial noise, consistent with people's fundamental ignorance of public health performance, despite the availability of official waiting time statistics in Spain.

Particularly germane to our analysis below, the figure makes clear that official regional waiting times explain a small amount of individual variation in beliefs. In particular, it turns out less

than 2% of individual belief variation is absorbed by the official waiting time.²² Similarly, partisanship and region together explain less than 6% of the variation.²³ In practice, this means that we have enough individual variation to test the prediction of the model: the four groups defined by partisanship crossed with assessment are fully populated in each of the regions.

5.2 The Structure of Beliefs

We next move to investigate how individuals attribute responsibility for the quality of public healthcare in Spain. We start by focusing on individuals who do not receive any exogenous information shock to their priors, that is, the 1274 individuals who were randomly assigned to the control group in our experiment.

5.2.1 Descriptive patterns

In Panel A of Figure 2, we plot the distribution of respondents' beliefs about the legal attribution of responsibility for healthcare. More specifically, the question asked is "Which public administration has authority over public healthcare provision?" and respondents are given five options: municipality, region, central government, European Union or "I don't know." The Figure shows a high degree of knowledge over the institutional allocation of responsibility, with approximately 60% of respondents who correctly identify that regions are *legally* responsible for managing healthcare services in Spain. Among the remaining respondents, 13% expressed uncertainty, and only 28% mistakenly attributed competence to the wrong government level. In short, citizens have a reasonably accurate grasp of the *de jure* allocation of responsibility.

In Panel B, we plot the distribution of respondents' beliefs about political responsibility for the quality of healthcare. As noted, the question asked is "Which layer of government do you believe is most responsible for the quality of public healthcare services?" and respondents choose a number on a scale from -10 to 10 where -10 indicates "all responsibility lies with the central government" and 10 indicates "all responsibility lies with the regional government." Here, a more

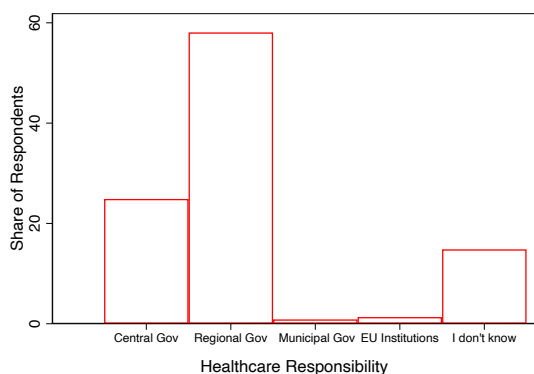
²²The R^2 of a regression predicting control group individual beliefs with regional official waiting times is 0.014. Details available from the authors.

²³The R^2 of a regression predicting control group individual beliefs with region-party fixed effects is 0.056. Details available from the authors.

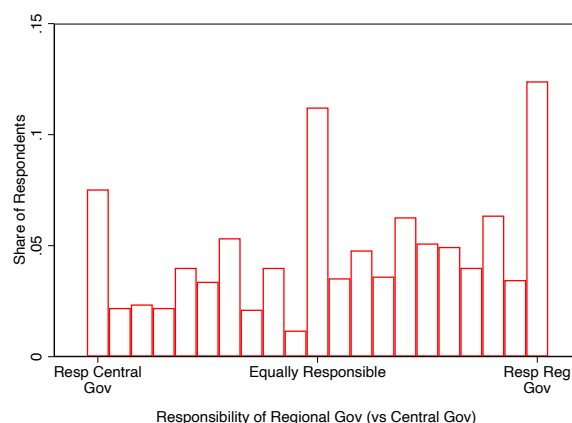
complex picture emerges: only 36% of respondents give a value of 5 or higher, which denotes most responsibility lies with the regional government. 22% answer with a value of -5 or lower, indicating Central Government responsibility and 11% viewed the central and regional governments as exactly equally responsible. Two points are worth noting here.

Figure 2: Distribution of Attribution of Responsibility

Panel A: Legal Responsibility



Panel B: Political Responsibility



Notes: Panel A displays the distribution of respondents' beliefs about legal (de jure) responsibility for healthcare management across levels of government. Panel B displays the distribution of respondents' beliefs about political responsibility for healthcare quality, measured on a scale from -10 (central government fully responsible) to $+10$ (regional government fully responsible). Sample restricted to control group respondents ($N = 1,274$).

First, respondents seem to have a different understanding of the concept of "legal authority" and the concept of "actual responsibility." This makes sense, as the legal framework in Spain is unambiguous in assigning healthcare to the regions. However, the fact that funding and framework laws within which this authority is exercised are decided by the Central government injects unavoidable ambiguity over who should be held accountable for failures. In any case, legal attribution is not treated by respondents as the ground truth of responsibility. This difference makes it hard to interpret studies of institutional ambiguity that inform respondents about which level of government is responsible.²⁴

²⁴Several studies experimentally manipulate responsibility attribution, for example by providing institutional responsibility cues or correcting respondents' beliefs about which level of government is formally responsible; see

Second, the distribution of beliefs over attribution displays a huge amount of variation. The standard deviation of this variable among respondents in the control group is 6.29 when the variable is capped by construction on a range of 21 units. There is therefore a lot of noise among people’s appreciation of responsibility despite the fact that all citizens live under uniform institutional rules.

These descriptive patterns already highlight the central empirical phenomenon our model is built to explain: despite reasonable knowledge of formal rules, there is substantial heterogeneity in how citizens attribute responsibility for performance across levels of government. This ambiguity is the backdrop against which partisan priors about competence can operate.

5.2.2 Testing the cross-sectional hypotheses

We now turn to testing our hypotheses $H_1 - H_5$, derived from Corollary 1. Recall that Table 1 summarizes the predictions. The key objects are (i) respondents’ ideology or partisanship (left versus right) and (ii) whether their priors classify them as having a good or bad assessment about waiting times.

Before presenting the formal tests, Figure 3 provides a graphical representation of the patterns we expect to observe in the population.²⁵ The figure plots respondents’ attribution of responsibility to the regional government against the difference between official regional waiting times and respondents’ prior beliefs. Therefore, as we move along the horizontal axis, respondents have a better assessment of healthcare performance. Panel A shows the relationship between assessment and attribution in divided regions. This relationship differs sharply by partisanship: among left-wing respondents, the fitted line slopes downward (i.e. the better they believe healthcare is, the *lower* their belief that the regional government is responsible for healthcare) whereas among right-wing respondents it slopes upward (i.e. the better they believe healthcare is, the *higher* their belief that the regional government is responsible for healthcare). This opposite partisan pattern

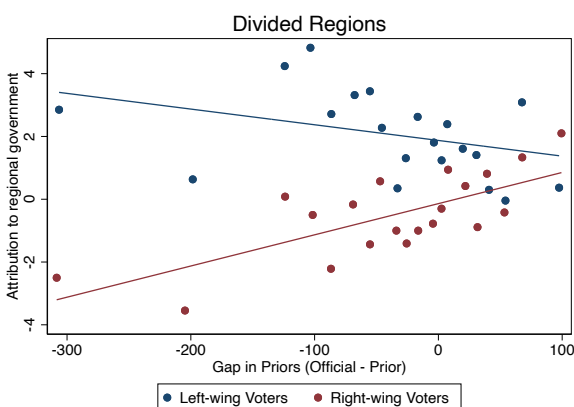
Tilley and Hobolt (2011), Malhotra and Kuo (2008), León and Orriols (2019), Fernández-Albertos et al. (2013), and Capezzone et al. (2026). Our evidence suggests that, in multilayer settings, such corrections should not be interpreted as revealing a unique ground truth of political accountability, since citizens may distinguish formal legal authority from actual responsibility for outcomes.

²⁵To limit the impact of extreme belief reports, in all figures the variables are winsorized at the 1st and 99th percentiles.

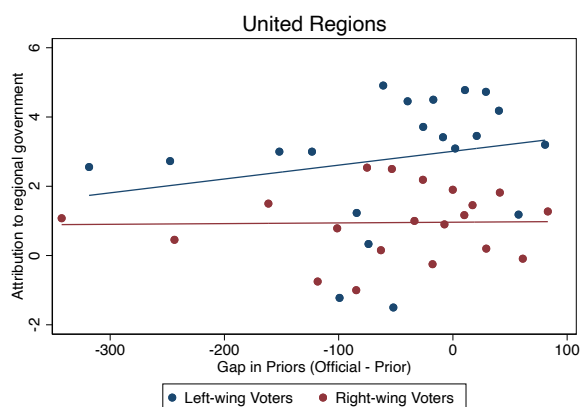
of attribution is precisely what the model predicts when responsibility is politically ambiguous. Panel B shows the corresponding pattern in united regions, where the same political bloc controls both levels of government. In this case, the crossed pattern is absent: both lines are flatter and, if anything, left-wing respondents attribute more responsibility to the regional government as their assessment of quality increases. The contrast between the two panels therefore previews the formal tests below. The empirical prediction is not simply that performance beliefs should correlate with attribution, but that this relationship should be structured by partisan priors only when institutional responsibility is politically ambiguous.

Figure 3: Performance Beliefs and Responsibility Attribution

Panel A: Divided Regions



Panel B: United Regions



Notes: The figure plots respondents' attribution of responsibility to the regional government against the difference between official regional waiting times and respondents' prior beliefs using binned scatter-plots. The vertical axis is respondents' belief about political responsibility for healthcare quality, measured on a scale from -10 to 10 , where higher values indicate greater responsibility attributed to the regional government. The horizontal axis is official waiting time minus the respondent's prior belief. Positive values indicate that official waiting times are longer than the respondent believed, while negative values indicate that official waiting times are shorter than the respondent believed. The sample is restricted to control group respondents. Linear fits are shown separately for left-wing and right-wing respondents. A region is classified as "divided" if none of the parties in the central government coalition participates in the regional government coalition, and as "united" otherwise; see Table C.1 for the classification of regions.

Turning to the formal tests, Table 3 shows the results in compact form as in Table 1. We dis-

play the statistic, as well as the corresponding one-sided p-value for the null that the coefficient is zero or has the sign opposite to the theoretical prediction. These statistics are derived from the estimation of equation (4) with data from divided regions –for the full regression details, please see column (1) of Table C.4 in the Appendix. Results for H_1 , H_2 and H_3 confirm the model’s predictions, all having the expected sign and being statistically significant at the 1% level (H_3) or at the 5% level (H_1 and H_2). For example, the coefficient on H_1 indicates that, on a scale from -10 to 10, left-wing good-assessment individuals attribute 1.09 points less responsibility to the regional government than left-wing bad-assessment individuals. The coefficient for H_3 is much larger, indicating that among bad-assessment, right-wing individuals attribute almost 4 points more responsibility to the central government than left-wing individuals. Recalling that the standard deviation of attribution is 6.29, these are economically as well as statistically significant results.

Table 3

Results in Divided Regions: Observational Evidence

	Left-wing	Right-wing	
Bad Assessment	Regional (+)	Central (–)	–3.83*** (0.00)
Good Assessment	Central (–)	Regional (+)	–1.05 (0.94)
	–1.09** (0.03)	1.69** (0.01)	2.78*** (0.00)

Notes: Entries report the relevant combinations of coefficient estimates testing hypotheses H1–H5 as defined in Table 1, with one-sided p -values in parentheses. Appendix Table C.4 reports the full set of estimated regression coefficients and standard errors. The model is estimated using equation (4). The sample is restricted to control group respondents in divided regions ($N = 819$). A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of divided regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Bad Assessment” and “Good Assessment” indicate whether respondents’ prior beliefs about waiting times exceed or fall below the official regional average, respectively. “Regional” means responsibility is attributed to the regional government, while “Central” means responsibility is attributed to the central government. All specifications include region fixed effects. Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

The coefficient on H_4 is non-significant and has a negative sign, which runs contrary to our hypothesis. This result would suggest that, if anything, right-wing good-assessment individuals

attribute less responsibility to the regional government than left-wing good-assessment individuals. One possible explanation for this anomaly is that right-wing individuals have a stronger normative belief on the role of the central government which may spill over on their factual beliefs. We have an inkling of this in Figure 3, as for both panels the line for Right-wing is systematically below the line for Left-wing voters. If such a partisan bias is present, as discussed in Section 4, the coefficient for H_3 should also be biased in the same direction, but the coefficient for H_5 , which is the difference between H_4 and H_3 should be correct. Indeed, Table 3 shows that H_5 is positive as predicted and statistically significant at the 1% level.

Overall, these descriptive results already display a clear pattern: attribution of responsibility in divided regions is strongly structured by the interaction of ideology and perceived performance in the way the theory predicts. These relationships are not mechanical consequences of formal responsibility or simple region-level confounders; they emerge after conditioning on region fixed effects and a rich set of controls.

To further bolster our interpretation of these results as supporting the theory, we repeat the same analysis in “united” regions, where the same political bloc controls both the central and regional governments. In these regions, the partisan updating we emphasize in the model cannot be present regardless of how responsibility is allocated across levels. They thus constitute valid placebos for our theory. Table 4 shows the results—for the full version, see column (2) of Appendix Table C.4. Consistent with the model, the partisan-performance structure in attribution is much weaker or entirely absent in united regions, as foreshadowed in Figure 3, Panel B: four of the five coefficients are insignificant and most have the wrong sign. Furthermore, the results suggest that, also in this subsample, right individuals systematically attribute more responsibility to the central government than left individuals.²⁶ This is consistent with $\zeta_{\mathcal{R}} < 0$ and suggests that the test statistics for both H_3 and H_4 are biased downwards in Table 3.

Finally, we use united regions as a control group, estimating the triple difference model given by equation (5). As discussed in Section 4, this allows us to estimate the hypotheses $H_1 - -H_5$ allowing for group-level confounders. Table 5 shows the results—for the full version, see column (3) of Appendix Table C.4. All the coefficients have signs consistent with our predictions, with the differences between divided and united regions in the key interaction terms precisely esti-

²⁶The coefficients are quantitatively sizable, at (-1.77) and (-2.36).

Table 4

Results in United Regions: Observational Evidence

	Left-wing	Right-wing	
Bad Assessment	?	?	−1.77** (0.01)
Good Assessment	?	?	−2.36 (0.99)
	0.98 (0.89)	0.39 (0.33)	−0.60 (0.69)

Notes: Entries report the relevant combinations of coefficient estimates testing the same H1–H5 hypotheses as in Table 3, with one-sided p -values in parentheses. Appendix Table C.4 reports the full set of estimated regression coefficients and standard errors. The model is estimated using equation (4). The sample is restricted to control group respondents in united regions ($N = 434$). A region is classified as “united” if at least one of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of united regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Bad Assessment” and “Good Assessment” indicate whether respondents’ prior beliefs about waiting times exceed or fall below the official regional average, respectively. All specifications include region fixed effects. Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

mated and matching the sign restrictions implied by the theory. In particular, for H_1 and H_3 , the coefficients are statistically significant at the 5% level, and for H_5 the coefficient is significant at the 1% level. The coefficients for H_2 and H_4 fall slightly short of significance at conventional levels (p-value of 0.12) but the magnitudes are sizable (1.30 and 1.31, respectively).

Table 5

Divided Regions Controlling for United: Observational Evidence

	Left-wing	Right-wing	
Bad Assessment	Regional (+)	Central (−)	−2.06** (0.01)
Good Assessment	Central (−)	Regional (+)	1.31 (0.12)
	−2.07** (0.02)	1.30 (0.12)	3.37*** (0.01)

Notes: Entries report the relevant combinations of coefficient estimates testing the same H1–H5 hypotheses as in Table 3, with one-sided p -values in parentheses. Appendix Table C.4 reports the full set of estimated regression coefficients and standard errors. The model is estimated using equation (5). $N = 1,253$. A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition, and as “united” otherwise; see Table C.1 for the classification of regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Bad Assessment” and “Good Assessment” indicate whether respondents’ prior beliefs about waiting times exceed or fall below the official regional average, respectively. “Regional” means responsibility is attributed to the regional government, while “Central” means responsibility is attributed to the central government. All specifications include region fixed effects. Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

This comparison strengthens the interpretation that the patterns we uncover in divided regions arise from attribution ambiguity interacting with partisan priors, rather than from broader ideological differences in how citizens interpret any political information. The magnitude of these differences, ranging between 1.3 and 2 points is remarkable against a standard deviation of attribution of 6.29.

One possible threat to interpreting these correlations is the demographic balance across groups. For example, if for some reason good-assessment individuals tended to be more educated, then results could simply reflect this omitted variable. In Section 5.4 we show that groups are balanced across all demographic characteristics we measure.

Finally, it is worth noting that in Appendix Table C.5 we present the results of running specification (4) and (5) using *de jure* responsibility as a dependent variable. No differences emerge as significant with the exception of the aforementioned fact that right-wing citizens systematically assign more responsibility to the central government. This lends additional support to the idea that citizens indeed think of legal and actual responsibilities as different objects, and that our model rightly operates only in the actual domain.

5.3 Informational Treatment

We now turn our attention to the informational treatment. As explained in Section 4, our experiment allows us to shock priors exogenously and thus test directly theoretical predictions derived from Propositions 1 and 4. The experiment also serves another purpose: while our theoretical approach derives beliefs on attribution from partisanship and beliefs on performance, the causal chain could plausibly go in the opposite direction. In a “partisan cheerleading” framework, for example, respondents answer performance questions as a function of their partisanship and beliefs over attribution.²⁷ For example, a right-wing voter who believes healthcare is the purview of the regional government ruled by Right, may as a consequence say that the waiting lists are short as a way of showing partisan support. This mechanism would generate correlations similar to what we predict and show in the cross-sectional results. By exogenously shocking beliefs on performance, we can ensure that in the causal chain attribution beliefs are an outcome and not a cause.

Learning and Belief Updating We begin by studying how respondents update their beliefs about waiting times when exposed to accurate, region-specific information. Recall that each respondent is asked for two waiting times (either specialist consultations or elective surgery) sequentially. Treated respondents are informed about the official regional waiting time relevant to the question they were first asked prior to being asked about the other waiting times. More specifically, half of treated respondents are asked first about their belief over the waiting time to visit with a medical specialist, then they are told the official regional average waiting time to visit with a medical specialist, and afterwards they are asked about their belief over the waiting

²⁷See Bullock et al. (2015).

time to undergo elective surgery. The other half of treated respondents are subject to the opposite sequence: asked about surgery, told about official waiting time for surgery, and then asked about specialist visit. This allows us to have a prior over waiting times (the answer to the first question) and a posterior over waiting times (the answer to the second question) under the assumption that citizens believe that there is a common component of healthcare quality driving both.

In the control group, we asked beliefs over both waiting times as well, sequentially. The difference with the treatment group is that we did not reveal any additional information between the two questions. In what follows, we treat the first answer as a prior and second as a posterior even though there is no treatment between the two.

Figure 4 displays the first-stage relationship graphically for treated respondents. The horizontal axis is the gap between the official regional waiting time and the respondent's prior belief, so positive values correspond to bad news and negative values to good news. The vertical axis is belief updating, defined as the posterior minus the prior. In all four panels, the slope is positive and close to the 45-degree line: respondents who learn that waiting times are longer than they expected revise their beliefs upward, while respondents who learn that waiting times are shorter than they expected revise their beliefs downward. Those who were farther from the official number in either direction, update more.²⁸ The pattern is very similar across united and divided regions and across left- and right-wing respondents. This indicates that the information treatment is understood, credible, and salient, and that respondents update their beliefs about healthcare performance in the expected direction.

Table 6 quantifies this learning using the standard learning-rate approach in information-provision experiments (Haaland et al., 2023).²⁹ The control group here is important for interpretation. In divided regions, control respondents exhibit a positive relationship between the official-prior gap and subsequent belief revision, with a coefficient of 0.49. Since these respondents receive no information between the two belief elicitations, this coefficient should not be interpreted as learning from treatment. Rather, it is a placebo updating benchmark. It may reflect

²⁸These patterns are all consistent with classical Bayesian learning.

²⁹The dependent variable is belief updating, defined as the posterior belief minus the prior belief, and the key explanatory variable is the gap between the official regional waiting time and the respondent's prior belief. A slope of one would correspond to full pass-through from the signal to the posterior belief, while a slope of zero would indicate no updating.

several forces: respondents may think more carefully about healthcare waiting times after the first elicitation, correct an initial typo or noisy response, display mechanical consistency across the two waiting-time measures, or report beliefs that share a common perceived component of regional healthcare quality across specialist consultations and elective surgery. The treatment effect should therefore be read as the additional movement toward the official waiting time relative to this control-group benchmark. In divided regions, treated respondents have a raw learning slope of 0.84. Netting out the baseline relationship observed in the control group, the treatment causes respondents to rationally update their beliefs by an additional 35 percentage points.

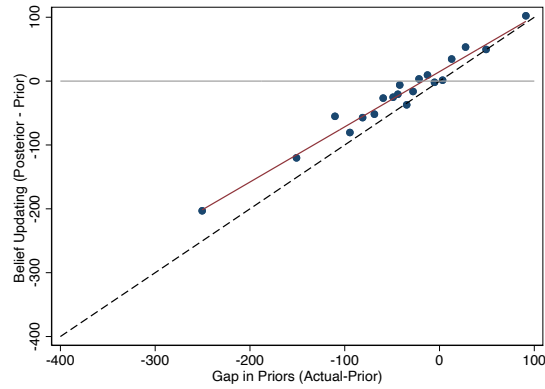
The same first-stage pattern appears in united regions. In the treatment-versus-control specification, the control-group slope is 0.41 and the coefficient on $\text{Treatment} \times \text{Bad-news gap}$ is 0.38, implying a treated-group slope of approximately 0.79. Thus, the magnitude of learning is very similar in divided and united regions. Importantly, there is little evidence that this first-stage updating differs by partisanship. The interaction between Right-wing respondent and the bad-news gap is small in both the control and treated split-sample specifications, and the triple interaction $\text{Right-wing respondent} \times \text{Treatment} \times \text{Bad-news gap}$ is close to zero in both divided and united regions. In other words, respondents update strongly on performance information, but this updating about waiting times themselves is not meaningfully partisan.

These magnitudes are well within the range of learning rates documented in prior information-provision experiments. The conservative net learning rates of 0.35–0.38 are close to existing estimates such as those for house prices and inflation expectations (Cavallo et al., 2017; Fuster et al., 2022b), while the raw treated slopes are at the high end of the literature and comparable to estimates from settings with highly salient or tailored quantitative information (Bottan and Perez-Truglia, 2022). This is plausible in our context because the treatment provides official, region-specific, numerical information about a concrete public-service outcome.

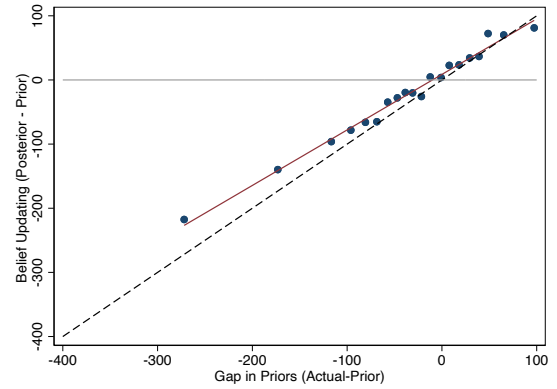
Overall, these results support the model’s premise that citizens can process precise performance information. Respondents are not simply stuck with their priors, nor do they refuse to revise beliefs when the signal is politically inconvenient. The key question for the theory is therefore not whether citizens learn about performance, but how this performance learning maps into attribution of responsibility when responsibility is ambiguous.

Figure 4: Learning and Belief Updating upon Treatment

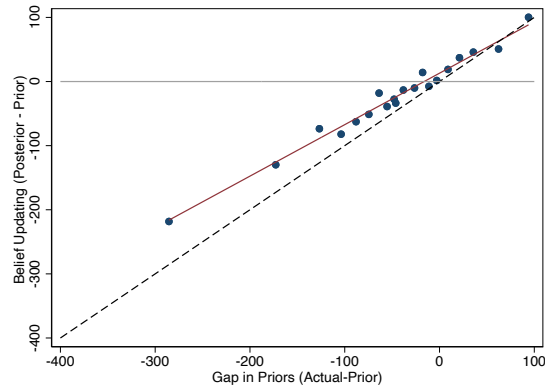
(a) United Regions, Left-Wing Voters



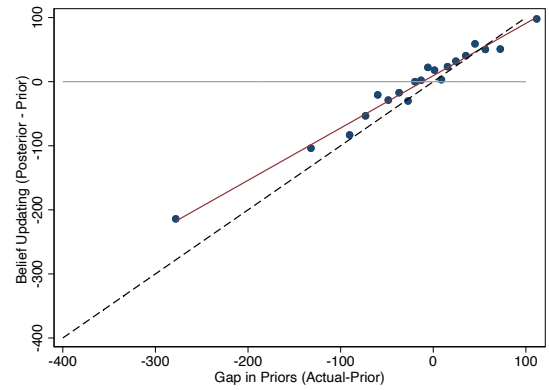
(b) Divided Regions, Left-Wing Voters



(c) United Regions, Right-Wing Voters



(d) Divided Regions, Right-Wing Voters



Notes: Binned scatterplots of belief updating (posterior minus prior) against the gap between official and prior waiting times, separately by region type (united vs. divided) and voter ideology (left-wing vs. right-wing). The solid red line shows the linear fit; the dashed 45-degree line indicates perfect updating to the official waiting time. A slope of one would indicate that respondents fully incorporate the new information. Sample restricted to treated respondents ($N = 2,583$). A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition, and as “united” otherwise; see Table C.1 for the classification of regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement.

Table 6

Learning and Belief Updating upon Treatment

	Control, Divided	Treated, Divided	Treatment vs. Control, Divided	Treatment vs. Control, United
	(1)	(2)	(3)	(4)
Right-wing Voter	-5.51 (4.73)	1.25 (3.33)	-5.62 (4.58)	-11.04** (6.61)
Bad-news gap (Official – Prior)	0.49*** (0.05)	0.84*** (0.03)	0.49*** (0.05)	0.41*** (0.06)
Right-wing Voter × Bad-news gap	-0.04 (0.08)	-0.03 (0.06)	-0.04 (0.08)	-0.00 (0.09)
Treatment			3.51 (3.94)	-6.97 (6.12)
Right-wing Voter × Treatment			6.88 (5.59)	15.10** (8.19)
Treatment × Bad-news gap			0.35*** (0.06)	0.38*** (0.08)
Right-wing Voter × Treatment × Bad-news gap			0.01 (0.10)	0.02 (0.12)
Constant	4.67* (3.38)	7.87*** (2.01)	4.48* (3.36)	17.17*** (5.01)
Observations	819	1666	2485	1321
R^2	0.29	0.57	0.50	0.48

Notes: The dependent variable is the difference between posterior and prior beliefs about waiting times (in days). Column 1 reports estimates for control group respondents in divided regions. Column 2 reports estimates for treated respondents in divided regions. Column 3 reports treatment vs. control estimates for all respondents in divided regions. Column 4 reports treatment vs. control estimates for all respondents in united regions. A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition, and as “united” otherwise; see Table C.1 for the classification of regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Bad Assessment” and “Good Assessment” indicate whether respondents’ prior beliefs about waiting times exceed or fall below the official regional average, respectively. All specifications include region fixed effects. Robust standard errors in parentheses. Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Finally, we use united regions to test for directional motives in belief updating. In these regions, the political valence of the signal is clear: since both layers of government are run by Left, bad news about waiting times is congenial to right-wing respondents, while good news is congenial to left-wing respondents. These irrational biases therefore predict asymmetric learning about performance itself: left-wing respondents should update more strongly when they receive good news than when they receive bad news, whereas right-wing respondents should update more strongly when they receive bad news than when they receive good news (e.g., Taber and Lodge, 2006). Appendix Table C.6 provides little evidence of this pattern. The treated-only estimates are directionally consistent with the hypothesis, but the better-identified treatment-versus-control specifications show no statistically significant treatment-induced asymmetry.

Before discussing our experimental results, it is worth noting that treatment randomization

balance was successful across demographic characteristics, as we show in Section 5.4.

Testing the treatment-effect hypotheses on attribution The core predictions of the model concern how new information about outcomes affects attribution. Proposition 1 argues that when citizens receive good news – learning that waiting times are shorter than they believed – the model predicts that they will rationally infer a higher probability that the actor they regard as more competent (i.e. their preferred party) was responsible; bad news has the opposite effect. As discussed in Section 4 this generates four hypotheses $HT_1 - HT_4$. Proposition 4 derives additional results about the directional differences in updating as a result of information revelation. For example, the gap in attribution to the regional government between right-wing and left-wing voters should expand or contract as a result of treatment. As discussed in Section 4, this logic generates five more hypotheses $HE_1 - HE_5$.

We begin by restricting the sample to individuals in divided regions and estimating equation (6). The relevant coefficients as well as p-values for the theoretical sign restrictions are shown in Table 7—for the full regression from which these estimates are derived, see column (1) of Appendix Table C.7. The results show that the effects of revealing the same performance information on attribution are strongly heterogeneous and mostly consistent with the theory. In particular, eight of the nine coefficients have the expected sign, with one being statistically significant at the 1% level, three being significant at the 5% level, one being significant at the 10% level, and three non-significant. When the information constitutes good news relative to respondents’ priors, right-wing bad-assessment individuals shift responsibility toward the regional government (their preferred side in our coding). When the information constitutes bad news, the signs reverse: right-wing good-assessment individuals reduce the responsibility they attribute to the regional government, while left-leaning good-assessment individuals increase it. In each case, citizens appear to be using the new performance information to update their beliefs about who is in charge in a way that favors the party they consider more competent, exactly as in the theoretical model. The only coefficient with a sign inconsistent with our hypotheses is the one corresponding to HT_1 , that is, the response of bad-assessment left-wing voters to good news.

Recalling that treatment closes the gap between the official and prior waiting times by about 36% compared to control, the size of some of these coefficients is economically significant. It

implies, for example, that moving a bad-assessment right-wing voter from priors to the official waiting time would shift responsibility towards the regional government by 2.11 points. Moving a good-assessment right-wing voter fully to the official waiting time would instead shift responsibility to the central government by 3.10 points.

Table 7

Treatment Effects in Divided Regions: Experimental Results

	Left-wing	Right-wing	
Good News	0.45 (0.85)	0.76* (0.08)	0.31 (0.33)
Bad News	0.45 (0.20)	−1.12** (0.04)	−1.57** (0.03)
	0.01 (0.50)	−1.87*** (0.01)	−1.88** (0.04)

Notes: Entries report the relevant combinations of coefficient estimates testing hypotheses HT1–HT4 and HE1–HE5 as defined in Table 2, with one-sided p -values in parentheses. Appendix Table C.7 reports the full set of estimated regression coefficients and standard errors. The model is estimated using equation (6). The sample includes treatment and control respondents in divided regions ($N = 2,485$). A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of divided regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Good News” indicates that the information treatment provides favorable news about waiting times relative to respondents’ prior beliefs, while “Bad News” indicates that it provides unfavorable news. All specifications include region fixed effects. Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

These patterns are not driven by generic shifts in attribution unrelated to performance. In Table 8 we show the placebo results when we restrict the sample to individuals in united regions—for the full results, see column (2) of Appendix Table C.7. Unlike the results for divided regions, the treatment effects in united regions provide no support for the predicted directional hypotheses, with estimates either statistically insignificant or falling in the opposite direction. This lends credibility to the theory, as it predicts no partisan effect on updating when there is no ambiguity on the party in charge.

Finally, we use united regions as a control group and estimate equation (8). The results are in Table 9—for the full results, see column (3) of Appendix Table C.7. Although results are more

Table 8

Placebo Effects in United Regions: Experimental Results

	Left-wing	Right-wing	
Good News	0.43	0.19	−0.24
	(0.76)	(0.38)	(0.61)
Bad News	−1.91	−0.79	1.13
	(0.99)	(0.17)	(0.85)
	−2.35	−0.98	1.37
	(0.99)	(0.17)	(0.83)

Notes: Entries report the relevant combinations of coefficient estimates testing the same HT1–HT4 and HE1–HE5 hypotheses as in Table 7, with one-sided p -values in parentheses. Appendix Table C.7 reports the full set of estimated regression coefficients and standard errors. The model is estimated using equation (6). The sample includes treatment and control respondents in united regions ($N = 1,321$). A region is classified as “united” if at least one of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of united regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Good News” indicates that the information treatment provides favorable news about waiting times relative to respondents’ prior beliefs, while “Bad News” indicates that it provides unfavorable news. In united regions, the same political bloc controls both government levels, so treatment effects on attribution mediated by partisan priors should be absent. All specifications include region fixed effects. Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

imprecise, eight of the nine coefficients have the expected sign, as before.³⁰

Table 9

Divided Regions Controlling for United: Experimental Results

	Left-wing	Right-wing	
Good News	0.01	0.56	0.55
	(0.51)	(0.25)	(0.31)
Bad News	2.37***	−0.33	−2.70**
	(0.01)	(0.37)	(0.03)
	2.35**	−0.90	−3.25**
	(0.03)	(0.25)	(0.03)

Notes: Entries report the relevant combinations of coefficient estimates testing the same HT1–HT4 and HE1–HE5 hypotheses as in Table 7, with one-sided p -values in parentheses. Appendix Table C.7 reports the full set of estimated regression coefficients and standard errors. The model is estimated using equation (8). $N = 3,806$. A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition, and as “united” otherwise; see Table C.1 for the classification of regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Good News” indicates that the information treatment provides favorable news about waiting times relative to respondents’ prior beliefs, while “Bad News” indicates that it provides unfavorable news. All specifications include region fixed effects. Significance levels: $*p < 0.10$, $**p < 0.05$, $***p < 0.01$.

Taken together, the experimental results provide a causal test of the mechanism highlighted by the theory. When we exogenously shock performance beliefs with precise information, citizens living under ambiguous partisan responsibility (i.e. in divided regions) update attribution in ways that are fully mediated by their partisan priors and the sign of the news.

It is again worth mentioning that in Appendix Table C.8 we run these experimental specifications using de jure responsibility as a dependent variable. Treatment effect coefficients are all very small and statistically insignificant, again consistent with the idea that upon learning new facts about performance, citizens update about actual responsibility but not about legal responsibility. Respondents in this experiment treat these two dimensions of responsibility very differently.

³⁰The coefficient for HT_1 still has the unexpected sign but it is now negligible in magnitude (0.01).

5.4 Balance and Robustness

Tables C.9–C.12 test for the balance in specifications (4), (5), (6), and (8), respectively. The tables show placebo estimates where the dependent variable in each specification is a socio-demographic variable. Hence, we should see zero effects on these variables. The results are reassuring, with no more significant effects than those expected by chance.

In Appendix Tables C.13–C.16 we study the robustness of the results using our four main specifications given by equations (4), (5), (6), and (8), respectively. In all tables, column (1) shows the baseline results, for reference. Column (2) drops the region fixed effects. Column (3) controls for socio-demographic variables. Columns (4)–(5) use alternative measures of right-wing partisanship. Column (4) defines an individual as “right” based on their ideology (ignoring past vote). Column (5) defines it based on past vote (ignoring ideology). Columns (6)–(9) use alternative measures of assessment. Column (6) defines good assessment as being below the median prior, regardless of official waiting time. Column (7) defines an individual as having a good assessment if she is below the median in the difference between the prior and ideal waiting time. Column (8) defines good assessment as being below the median prior, where the median prior is defined separately for each combination of region and individual’s ideology. Column (9) considers a continuous variable, defined as the standardized difference in Official - Prior. Column (10) drops Catalonia from the sample. The results are remarkably robust to all these changes.

Finally, Appendix Table C.17 reports results of specifications where the two experimental treatments are treated separately instead of combined into one. The results reveal similar effects as in our main results.

6 Discussion and Conclusion

This paper develops and tests a theory of belief updating under attribution ambiguity. We show that when voters face uncertainty about which level of government is responsible for observed policy outcomes, rational Bayesian updating generates patterns that have often been interpreted as evidence of motivated reasoning or partisan bias: citizens attribute good outcomes to their preferred party and bad outcomes to the party they dislike. The key insight is that performance signals are jointly informative about both the quality of outcomes and the identity of the re-

sponsible actor. When responsibility is ambiguous, voters with different priors about politicians' competence rationally draw different inferences about attribution from the same signal.

We formalize this logic in a model where citizens observe policy outcomes but not which of two actors produced them. In the model partisanship simply reflects a voter's belief over the relative competence of the two parties. The model yields sharp comparative statics: after good outcomes, voters attribute responsibility toward the party they believe to be more competent; after bad outcomes, they shift blame toward the opposing party. Crucially, these patterns persist even when voters receive precise information about performance, because such information does not resolve the underlying identification problem regarding attribution.

We test these predictions using original survey data from Spain, where both regional and central governments plausibly share responsibility for healthcare. The observational evidence confirms the model's cross-sectional predictions: in regions with divided government, good-assessment voters credit their preferred party while bad-assessment voters shift responsibility to the opposition. These patterns are absent in regions where the same political bloc controls both levels of government, so attribution does not carry the same partisan contrast. Our information experiment provides causal evidence consistent with the model: exogenous news about health-care quality induces heterogeneous updating of responsibility beliefs that varies systematically with both the direction of news and voters' partisan priors.

Three implications follow from our analysis. First, the literature's emphasis on distinguishing rational from biased updating may be less productive than previously thought. In environments with attribution ambiguity—which characterize most multilevel democracies—observationally identical updating patterns can arise under either assumption. Second, information campaigns about policy performance may have limited effects on accountability not because voters are irrational, but because such information does not resolve the deeper uncertainty about who deserves credit or blame. Efforts to improve accountability may need to address institutional clarity alongside informational gaps. Third, our results suggest that the fragmentation of responsibility across levels of government creates structural obstacles to electoral accountability that cannot be overcome through voter education alone.

Beyond political settings, these results have implications for organizational economics. If there is ambiguity over who or why a particular project has succeeded or failed, our framework

suggests that priors will loom very heavily over the conclusions managers may draw and substantial disagreements may arise about the way forward even when information about what has transpired is shared and precise.

Future work might examine how politicians strategically manipulate attribution, how institutional reforms that clarify responsibility affect accountability, and whether repeated interactions allow voters to learn about the mapping from outcomes to responsibility. Our framework provides a foundation for analyzing these questions while highlighting that the challenge of democratic accountability in complex polities runs deeper than voter ignorance or bias.

References

- Acemoglu, Daron, Victor Chernozhukov, and Muhamet Yildiz**, “Fragility of Asymptotic Agreement under Bayesian Learning,” *Theoretical Economics*, 2016, 11 (1), 187–225.
- Adida, Claire, Jessica Gottlieb, Eric Kramon, and Gwyneth McClendon***, “Reducing or reinforcing in-group preferences? An experiment on information and ethnic voting,” *Quarterly Journal of Political Science*, 2017, 12 (4), 437–477.
- Anderson, Cameron D.**, “Economic Voting and Multilevel Governance: A Comparative Individual-Level Analysis,” *American Journal of Political Science*, April 2006, 50 (2), 449–463.
- Ashworth, Scott**, “Electoral accountability: Recent theoretical and empirical work,” *Annual Review of Political Science*, 2012, 15 (1), 183–201.
- , “Electoral Accountability: Recent Theoretical and Empirical Work,” *Annual Review of Political Science*, June 2012, 15 (1), 183–201.
- , **Ethan Bueno De Mesquita, and Amanda Friedenberg**, “Accountability and Information in Elections,” *American Economic Journal: Microeconomics*, May 2017, 9 (2), 95–138.
- Banerjee, Abhijit, Nils Enevoldsen, Rohini Pande, and Michael Walton**, “Public Information Is an Incentive for Politicians: Experimental Evidence from Delhi Elections,” *American Economic Journal: Applied Economics*, 2024, 16 (3), 323–353.
- Bartels, Larry M**, “Beyond the running tally: Partisan bias in political perceptions,” *Political behavior*, 2002, 24 (2), 117–150.
- Besley, Timothy and Andrea Prat**, “Handcuffs for the Grabbing Hand? Media Capture and Government Accountability,” *American Economic Review*, May 2006, 96 (3), 720–736.
- Bhandari, Abhit, Horacio Larreguy, and John Marshall**, “Able and Mostly Willing: An Empirical Anatomy of Information’s Effect on Voter-Driven Accountability in Senegal,” *American Journal of Political Science*, October 2023, 67 (4), 1040–1066.

- Bisbee, James and Diana Da In Lee**, “Objective facts and elite cues: partisan responses to covid-19,” *The Journal of Politics*, 2022, 84 (3), 1278–1291.
- Bisgaard, Martin**, “How getting the facts right can fuel partisan-motivated reasoning,” *American Journal of Political Science*, 2019, 63 (4), 824–839.
- Bonomi, Giampaolo, Nicola Gennaioli, and Guido Tabellini**, “Identity, Beliefs, and Political Conflict,” *The Quarterly Journal of Economics*, October 2021, 136 (4), 2371–2411.
- Bottan, Nicolas L and Ricardo Perez-Truglia**, “Choosing your pond: location choices and relative income,” *Review of Economics and Statistics*, 2022, 104 (5), 1010–1027.
- Brockmeyer, Anne, Francisco Garfias, and Juan Carlos Suárez Serrato**, “The Fiscal Contract up Close: Experimental Evidence from Mexico City,” Working Paper 32776, National Bureau of Economic Research 2024.
- Bullock, John G, Alan S Gerber, Seth J Hill, and Gregory A Huber**, “Partisan bias in factual beliefs about politics,” *Quarterly Journal of Political Science*, 2015, 10 (4), 519–578.
- Capezzone, Tommaso, Pierluigi Conzo, Willem Sas, Dmitriy Vorobyev, and Roberto Zotti**, “Who Is to Blame (or Praise)? Information, Responsibility Attribution, and Service Quality in Multilevel Government,” Technical Report, University of Turin 2026.
- Cavallo, Alberto, Guillermo Cruces, and Ricardo Perez-Truglia**, “Inflation expectations, learning, and supermarket prices: Evidence from survey experiments,” *American Economic Journal: Macroeconomics*, 2017, 9 (3), 1–35.
- Centro de Investigaciones Sociológicas (CIS)**, “Casi el 80% de la población utilizó los servicios de Atención Primaria en 2024,” February 2025. Sala de prensa.
- Chong, Alberto, Ana L. De La O, Dean Karlan, and Leonard Wantchekon**, “Does Corruption Information Inspire the Fight or Quash the Hope? A Field Experiment in Mexico on Voter Turnout, Choice, and Party Identification,” *The Journal of Politics*, 2015, 77 (1), 55–71.
- Diaconis, Persi and David Freedman**, “On the Uniform Consistency of Bayes Estimates for Multinomial Probabilities,” *The Annals of Statistics*, 1990, 18 (3), 1317–1327.

- Dunning, Thad, Guy Grossman, Macartan Humphreys, Susan D Hyde, Craig McIntosh, Gareth Nellis, Claire L Adida, Eric Arias, Clara Bicalho, Taylor C Boas et al.**, “Voter information campaigns and political accountability: Cumulative findings from a preregistered meta-analysis of coordinated trials,” *Science advances*, 2019, 5 (7), eaaw2612.
- Fearon, James D.**, “Electoral Accountability and the Control of Politicians: Selecting Good Types versus Sanctioning Poor Performance,” in Adam Przeworski, Susan C. Stokes, and Bernard Manin, eds., *Democracy, Accountability, and Representation*, Cambridge: Cambridge University Press, 1999, pp. 55–97.
- Fernández-Albertos, José, Alexander Kuo, and Laia Balcells**, “Economic Crisis, Globalization, and Partisan Bias: Evidence from Spain,” *International Studies Quarterly*, 2013, 57 (4), 804–816.
- Ferraz, Claudio and Frederico Finan**, “Exposing corrupt politicians: the effects of Brazil’s publicly released audits on electoral outcomes,” *The Quarterly journal of economics*, 2008, 123 (2), 703–745.
- Fuster, Andreas, Ricardo Perez-Truglia, Mirko Wiederholt, and Basit Zafar**, “Expectations with Endogenous Information Acquisition: An Experimental Investigation,” *The Review of Economics and Statistics*, September 2022, 104 (5), 1059–1078.
- , **Ricardo Perez-Truglia, Mirko Wiederholt, and Basit Zafar**, “Expectations with endogenous information acquisition: An experimental investigation,” *Review of Economics and Statistics*, 2022, 104 (5), 1059–1078.
- Gentzkow, Matthew, Michael B. Wong, and Allen T. Zhang**, “Ideological Bias and Trust in Information Sources,” *American Economic Journal: Microeconomics*, 2025, 17 (2), 162–213.
- Graham, Matthew H. and Shikhar Singh**, “An Outbreak of Selective Attribution: Partisanship and Blame in the COVID-19 Pandemic,” *American Political Science Review*, February 2024, 118 (1), 423–441.

- Grossman, Guy and Kristin Michelitch**, “Information Dissemination, Competitive Pressure, and Politician Performance between Elections: A Field Experiment in Uganda,” *American Political Science Review*, 2018, 112 (2), 280–301.
- Haaland, Ingar, Christopher Roth, and Johannes Wohlfart**, “Designing information provision experiments,” *Journal of economic literature*, 2023, 61 (1), 3–40.
- Healy, Andrew, Alexander G Kuo, and Neil Malhotra**, “Partisan bias in blame attribution: when does it occur?,” *Journal of Experimental Political Science*, 2014, 1 (2), 144–158.
- Heinkelmann-Wild, Tim, Bernhard Zangl, Berthold Rittberger, and Lisa Kriegmair**, “Blame Shifting and Blame Obfuscation: The Blame Avoidance Effects of Delegation in the European Union,” *European Journal of Political Research*, February 2023, 62 (1), 221–238.
- Hobolt, Sara B. and James Tilley**, “Who’s in Charge? How Voters Attribute Responsibility in the European Union,” *Comparative Political Studies*, May 2014, 47 (6), 795–819.
- Hurst, Reuben, Andrew Simon, and Michael Ricks**, “Public Good Perceptions and Polarization: Evidence from Higher Education Appropriations,” *SSRN Electronic Journal*, 2024.
- Incerti, Trevor**, “Corruption Information and Vote Share: A Meta-Analysis and Lessons for Experimental Design,” *American Political Science Review*, August 2020, 114 (3), 761–774.
- León, Sandra**, “Who is responsible for what? Clarity of responsibilities in multilevel states: The case of Spain,” *European Journal of Political Research*, 2011, 50 (1), 80–109.
- **and Lluís Orriols**, “Attributing Responsibility in Devolved Contexts: Experimental Evidence from the UK,” *Electoral Studies*, 2019, 59, 39–48.
- **, Ignacio Jurado, and Amuitz Garmendia Madariaga**, “Passing the Buck? Responsibility Attribution and Cognitive Bias in Multilevel Democracies,” *West European Politics*, May 2018, 41 (3), 660–682.
- Little, Andrew T.**, “How to Distinguish Motivated Reasoning from Bayesian Updating,” *Political Behavior*, 2025, 47, 1501–1525.

- , **Keith E. Schnakenberg**, and **Ian R. Turner**, “Motivated Reasoning and Democratic Accountability,” *American Political Science Review*, 2022, 116 (2), 751–767.
- Malhotra, Neil and Alexander G. Kuo**, “Attributing Blame: The Public’s Response to Hurricane Katrina,” *The Journal of Politics*, 2008, 70 (1), 120–135.
- Martin, Lucy and Pia J. Raffler**, “Fault Lines: The Effects of Bureaucratic Power on Electoral Accountability,” *American Journal of Political Science*, January 2021, 65 (1), 210–224.
- Martinez-Bravo, Monica and Carlos Sanz**, “Trust and Accountability in Times of Crisis,” *Economica*, January 2025, 92 (365), 230–258.
- Powell, G. Bingham and Guy D. Whitten**, “A Cross-National Analysis of Economic Voting: Taking Account of the Political Context,” *American Journal of Political Science*, May 1993, 37 (2), 391.
- Stantcheva, Stefanie**, “Understanding Tax Policy: How Do People Reason?,” *The Quarterly Journal of Economics*, 2021, 136 (4), 2309–2369.
- Taber, Charles S and Milton Lodge**, “Motivated skepticism in the evaluation of political beliefs,” *American journal of political science*, 2006, 50 (3), 755–769.
- Tilley, James and Sara B Hobolt**, “Is the government to blame? An experimental test of how partisanship shapes perceptions of performance and responsibility,” *The journal of politics*, 2011, 73 (2), 316–330.

A Theoretical Appendix

Proof of Proposition 1. Consider a generic citizen (we drop the subscript i for simplicity) and recall that $\alpha(a)$ is the prior probability that $a = 1$. We start by deriving the (marginal) posteriors resulting from Bayesian updating after observing outcome $Y \in \{S, F\}$. After $Y = S$,

$$f_R(\theta_R|Y = S) = \frac{\Pr[Y = S|\theta_R]}{\Pr[Y = S]} f_R(\theta_R) = \frac{\alpha(1)\theta_R + \alpha(0)\mathbb{E}[\theta_L]}{\alpha(1)\mathbb{E}[\theta_R] + \alpha(0)\mathbb{E}[\theta_L]} f_R(\theta_R), \quad (9)$$

$$f_L(\theta_L|Y = S) = \frac{\alpha(1)\mathbb{E}[\theta_R] + \alpha(0)\theta_L}{\alpha(1)\mathbb{E}[\theta_R] + \alpha(0)\mathbb{E}[\theta_L]} f_L(\theta_L), \quad (10)$$

$$p(Y = S) = \Pr(a = 1|Y = S) = \frac{\alpha(1)\mathbb{E}[\theta_R]}{\alpha(1)\mathbb{E}[\theta_R] + \alpha(0)\mathbb{E}[\theta_L]}, \quad (11)$$

and similarly, if $Y = F$

$$f_R(\theta_R|Y = F) = \frac{\alpha(1)(1 - \theta_R) + \alpha(0)\mathbb{E}[1 - \theta_L]}{\alpha(1)\mathbb{E}[1 - \theta_R] + \alpha(0)\mathbb{E}[1 - \theta_L]} f_R(\theta_R), \quad (12)$$

$$f_L(\theta_L|Y = F) = \frac{\alpha(1)\mathbb{E}[1 - \theta_R] + \alpha(0)(1 - \theta_L)}{\alpha(1)\mathbb{E}[1 - \theta_R] + \alpha(0)\mathbb{E}[1 - \theta_L]} f_L(\theta_L), \quad (13)$$

$$p(Y = F) = \Pr(a = 1|Y = F) = \frac{\mathbb{E}[1 - \theta_R]\alpha(1)}{\alpha(1)\mathbb{E}[1 - \theta_R] + \alpha(0)\mathbb{E}[1 - \theta_L]}. \quad (14)$$

Let $\Delta_R(Y) = \mathbb{E}[\theta_R - \theta_L|Y]$ be R 's advantage over L given outcome Y and $\Delta_R(\emptyset) = \mathbb{E}[\theta_R - \theta_L]$ be the prior advantage. Then, the effect of a success on R 's advantage is

$$\begin{aligned} \Delta_R(Y = S) &= \frac{1}{\Pr[Y = S]} \left(\alpha(1) \left(\int_0^1 \theta_R^2 f_R(\theta_R) d\theta_R - \mathbb{E}[\theta_R]\mathbb{E}[\theta_L] \right) + \alpha(0) \left(\mathbb{E}[\theta_L]\mathbb{E}[\theta_R] - \int_0^1 \theta_L^2 f_L(\theta_L) d\theta_L \right) \right) \\ &= \frac{1}{\Pr[Y = S]} (\alpha(1) (Var[\theta_R] + \mathbb{E}[\theta_R] (\mathbb{E}[\theta_R] - \mathbb{E}[\theta_L])) - \alpha(0) (Var[\theta_L] + \mathbb{E}[\theta_L] (\mathbb{E}[\theta_L] - \mathbb{E}[\theta_R]))) \\ &= \frac{\alpha(1)Var[\theta_R] - \alpha(0)Var[\theta_L]}{\Pr[Y = S]} + \Delta_R(\emptyset). \end{aligned} \quad (15)$$

Similarly, we have

$$\begin{aligned} \Delta_R(Y = F) &= \frac{\alpha(1)}{\Pr[Y = F]} \left(\int (1 - \theta_R)\theta_R f_R(\theta_R) d\theta_R - \mathbb{E}[\theta_L]\mathbb{E}[1 - \theta_R] \right) \\ &\quad + \frac{\alpha(0)}{\Pr[Y = F]} \left(\mathbb{E}[\theta_R]\mathbb{E}[1 - \theta_L] - \int (1 - \theta_L)\theta_L f_L(\theta_L) d\theta_L \right) \\ &= \frac{\alpha(0)Var[\theta_L] - \alpha(1)Var[\theta_R]}{\Pr[Y = F]} + \Delta_R(\emptyset). \end{aligned} \quad (16)$$

Finally, (11) and (14) confirm that attribution satisfies

$$\frac{p(Y = S)}{1 - p(Y = S)} = \frac{\mathbb{E}[\theta_R]}{\mathbb{E}[\theta_L]} \frac{\alpha(1)}{1 - \alpha(1)} \text{ and } \frac{p(Y = F)}{1 - p(Y = F)} = \frac{\mathbb{E}[1 - \theta_R]}{\mathbb{E}[1 - \theta_L]} \frac{\alpha(1)}{1 - \alpha(1)}. \quad (17)$$

Therefore, $p(Y = S) > \alpha(1)$ iff $\mathbb{E}[\theta_R] > \mathbb{E}[\theta_L]$ —i.e., if R is voter's initial favorite politician—while $p(Y = F) > \alpha(1)$ iff $\mathbb{E}[\theta_L] > \mathbb{E}[\theta_R]$ —i.e., if L is voter's initial favorite politician. $\square$

Proof of Proposition 2. Using (15) and (16) we have for voter i

$$\tilde{\Delta}_R^i(Y_i = S) = \frac{\alpha(1)Var_i[\theta_R] - \alpha(0)Var_i[\theta_L]}{\Pr[Y_i = S]} \quad (18)$$

$$\tilde{\Delta}_R^i(Y_i = F) = \frac{\alpha(0)Var_i[\theta_L] - \alpha(1)Var_i[\theta_R]}{\Pr[Y_i = F]} \quad (19)$$

Thus, $\tilde{\Delta}_R^i(Y_i = S) > 0$ iff (1) holds, while $\tilde{\Delta}_R^i(Y_i = F) < 0$ iff (1) holds. In particular, setting $\alpha_i(1) = 1$ and $\alpha_i(0) = 0$ in (18) and (19) yields $\tilde{\Delta}_R(Y_i = S) > 0$ and $\tilde{\Delta}_R(Y_i = F) < 0$ for all voters. Conversely, setting $\alpha_i(0) = 1$ and $\alpha_i(1) = 0$ in (18) and (19) yields $\tilde{\Delta}_R(Y_i = S) < 0$ and $\tilde{\Delta}_R(Y_i = F) > 0$ for all voters. $\square$

To prove Proposition 3, we use the following Lemma.

Lemma 1. *Suppose that citizens have a uniform prior over attribution—i.e., for all i , $\alpha_i(a = 1) = \alpha_i(a = 0) = 1/2$. If $i \in \mathcal{R}$, then R remains voter i 's favorite politician if $Y_i = S$ and $(\mathbb{E}_i[\theta_R])^2 \geq \mathbb{E}_i[\theta_L]$, or if $Y = F$ and $(\mathbb{E}_i[1 - \theta_L])^2 \geq \mathbb{E}_i[1 - \theta_R]$. If instead $i \in \mathcal{L}$, then L remains voter i 's favorite politician if $Y = S$ and $(\mathbb{E}_i[\theta_L])^2 \geq \mathbb{E}_i[\theta_R]$, or if $Y = F$ and $(1 - \mathbb{E}_i[\theta_R])^2 \geq 1 - \mathbb{E}_i[\theta_L]$.*

Proof of Lemma 1. For this proof, we drop the subscript i for simplicity. Since $0 \leq \theta_j \leq 1$ then $\mathbb{E}[\theta_j^2] \leq \mathbb{E}[\theta_j]$, while $Var[\theta_j] \geq 0$ implies that $(\mathbb{E}[\theta_j])^2 \leq \mathbb{E}[\theta_j^2]$. These inequalities provide the bounds

$$(\mathbb{E}[\theta_R])^2 - \mathbb{E}[\theta_L] \leq \mathbb{E}[\theta_R^2] - \mathbb{E}[\theta_L^2] \leq \mathbb{E}[\theta_R] - (\mathbb{E}[\theta_L])^2, \quad (20)$$

$$(\mathbb{E}[1 - \theta_L])^2 - \mathbb{E}[1 - \theta_R] \leq \mathbb{E}[(1 - \theta_L)^2] - \mathbb{E}[(1 - \theta_R)^2] \leq \mathbb{E}[1 - \theta_L] - (\mathbb{E}[1 - \theta_R])^2. \quad (21)$$

We can simplify (15) and (16) using $\alpha(1) = \alpha(0) = 1/2$,

$$\begin{aligned} \Delta_R(Y = S) &= \frac{\alpha(1)Var[\theta_R] - \alpha(0)Var[\theta_L]}{\Pr[Y = S]} + \Delta_R(\emptyset) = \frac{\mathbb{E}[\theta_R^2] - \mathbb{E}[\theta_L^2]}{\mathbb{E}[\theta_R] + \mathbb{E}[\theta_L]}, \\ \Delta_R(Y = F) &= \frac{\alpha(0)Var[\theta_L] - \alpha(1)Var[\theta_R]}{\Pr[Y = F]} + \Delta_R(\emptyset) = \frac{\mathbb{E}[(1 - \theta_L)^2] - \mathbb{E}[(1 - \theta_R)^2]}{\mathbb{E}[1 - \theta_R] + \mathbb{E}[1 - \theta_L]}. \end{aligned}$$

Using the bounds (20) and (21) we then obtain the following bounds on $\Delta_R(Y)$

$$\begin{aligned} (\mathbb{E}[\theta_R])^2 - \mathbb{E}[\theta_L] &\leq \Delta_R(Y = S)(\mathbb{E}[\theta_R] + \mathbb{E}[\theta_L]) \leq \mathbb{E}[\theta_R] - (\mathbb{E}[\theta_L])^2, \\ (\mathbb{E}[1 - \theta_L])^2 - (\mathbb{E}[1 - \theta_R]) &\leq \Delta_R(Y = F)(\mathbb{E}[1 - \theta_R] + \mathbb{E}[1 - \theta_L]) \leq \mathbb{E}[1 - \theta_L] - (1 - \mathbb{E}[\theta_R])^2. \end{aligned}$$

Consider first the case of sufficiency. If $(\mathbb{E}[\theta_R])^2 \geq \mathbb{E}[\theta_L]$ —which implies $\mathbb{E}[\theta_R] \geq \mathbb{E}[\theta_L]$ so R is the voter's initial favorite politician—then $\Delta_R(Y = S) \geq 0$ regardless of the variance that the voter faces, while if $(\mathbb{E}[\theta_L])^2 \geq \mathbb{E}[\theta_R]$ —which implies $\mathbb{E}[\theta_L] > \mathbb{E}[\theta_R]$ so L is the voter's initial favorite politician—then $\Delta_R(Y = S) \leq 0$, again, regardless of the variance of θ_R and θ_L . Moreover, if $\mathbb{E}[\theta_R] \geq 1 - (1 - \mathbb{E}[\theta_L])^2$ —which implies $1 - \mathbb{E}[\theta_L] > 1 - \mathbb{E}[\theta_R]$ so R is the voter's initial favorite politician—then $\Delta_R(Y = F) > 0$ while if $\mathbb{E}[\theta_L] > 1 - (1 - \mathbb{E}[\theta_R])^2$ —which implies $1 - \mathbb{E}[\theta_L] < 1 - \mathbb{E}[\theta_R]$ so L is the voter's initial favorite politician—then $\Delta_R(Y = F) < 0$, in both cases regardless of the variance of θ_R and θ_L .

Consider now the case of necessity: we will find prior beliefs that reverse partisanship after the voter privately observes the policy outcome. Suppose first that $(\mathbb{E}[\theta_R])^2 < \mathbb{E}[\theta_L]$ but $\text{Var}[\theta_R] = \varepsilon$ and $\text{Var}[\theta_L] = \mathbb{E}[\theta_L](1 - \mathbb{E}[\theta_L]) - \varepsilon$ with $0 < 2\varepsilon < \mathbb{E}[\theta_L] - (\mathbb{E}[\theta_R])^2$. Then,

$$\Delta_R(Y = S) = \frac{\varepsilon + \varepsilon - \mathbb{E}[\theta_L](1 - \mathbb{E}[\theta_L]) + \mathbb{E}[\theta_R]^2 - \mathbb{E}[\theta_L]^2}{\mathbb{E}[\theta_R] + \mathbb{E}[\theta_L]} = \frac{2\varepsilon + \mathbb{E}[\theta_R]^2 - \mathbb{E}[\theta_L]}{\mathbb{E}[\theta_R] + \mathbb{E}[\theta_L]} < 0. \quad (22)$$

If instead $\mathbb{E}[1 - \theta_L]^2 < \mathbb{E}[1 - \theta_R]$, consider prior variances $\text{Var}[\theta_R] = \mathbb{E}[\theta_R](1 - \mathbb{E}[\theta_R]) - \varepsilon$ and $\text{Var}[\theta_L] = \varepsilon$ with $0 < 2\varepsilon < \mathbb{E}[1 - \theta_R] - \mathbb{E}[1 - \theta_L]^2$. Then,

$$\Delta_R(Y = F) = \frac{\varepsilon + \varepsilon - \mathbb{E}[\theta_R](1 - \mathbb{E}[\theta_R]) + \mathbb{E}[1 - \theta_L]^2 - \mathbb{E}[1 - \theta_R]^2}{\mathbb{E}[1 - \theta_R] + \mathbb{E}[1 - \theta_L]} \quad (23)$$

$$= \frac{2\varepsilon + \mathbb{E}[1 - \theta_L]^2 - \mathbb{E}[1 - \theta_R]}{\mathbb{E}[1 - \theta_R] + \mathbb{E}[1 - \theta_L]} < 0. \quad (24)$$

We can apply the same logic to the case of an $\mathcal{L}$ -voter to prove sufficiency in part ii. □

Proof of Proposition 3. Using (15) and (17), Bayesian updating implies that after observing $Y_i = S$, voter i will reach posteriors

$$p_i(Y_i = S) = \frac{\mathbb{E}_i[\theta_R]}{\mathbb{E}_i[\theta_R] + \mathbb{E}_i[\theta_L]} \text{ and } \Delta_R^i(Y_i = S) = \frac{\mathbb{E}_i[\theta_R^2] - \mathbb{E}_i[\theta_L^2]}{\mathbb{E}_i[\theta_R] + \mathbb{E}_i[\theta_L]}, \quad (25)$$

while, using (16) and (17) we have,

$$p_i(Y_i = F) = \frac{\mathbb{E}_i[1 - \theta_R]}{\mathbb{E}_i[1 - \theta_R] + \mathbb{E}_i[1 - \theta_L]} \text{ and } \Delta_R^i(Y_i = F) = \frac{\mathbb{E}_i[(1 - \theta_L)^2] - \mathbb{E}_i[(1 - \theta_R)^2]}{\mathbb{E}_i[1 - \theta_R] + \mathbb{E}_i[1 - \theta_L]}, \quad (26)$$

after $Y_i = F$. This implies that $p_i(Y_i = S) \geq 1/2$ iff $\mathbb{E}_i[\theta_R] \geq \mathbb{E}_i[\theta_L]$, while $p_i(Y_i = F) \geq 1/2$ iff $\mathbb{E}_i[1 - \theta_R] \geq \mathbb{E}_i[1 - \theta_L]$, so that attribution follows initial partisanship: for voter i , $A_R(Y = S) = 1$ and $A_L(Y = F) = 1$ if $i \in \mathcal{R}$, while if $i \in \mathcal{L}$, then $A_L(Y = S) = 1$ and $A_R(Y = F) = 1$.

Lemma 1 provides the maximum mass of voters that may switch allegiance after observing their private outcome. Indeed, Lemma 1 shows that all $i \in \mathcal{R}$ with $(\mathbb{E}_i[\theta_R])^2 \geq \mathbb{E}_i[\theta_L]$ that experience $Y_i = S$, or experience $Y_i = F$ and $(1 - \mathbb{E}_i[\theta_L])^2 \geq 1 - \mathbb{E}_i[\theta_R]$, are guaranteed to still see R as their favorite after observing the policy outcome; thus, these inequalities indeed define $base_R(Y)$. Lemma 1 also shows that any $\mathcal{R}$ -voter that is not in $base_R(Y)$ may potentially switch allegiance; that is, there are prior beliefs for any voter in $crossover_R(Y) = \mathcal{R} \cap (base_R(Y))^c$ that would lead the voter to switch allegiance after observing Y , and similarly for any voter in $crossover_L(Y) = \mathcal{L} \cap (base_L(Y))^c$.

From the definition of $\rho_{jk}(Y)$, and focusing on $j = k = R$ we can write

$$\frac{\rho_{RR}(Y)}{1 - \rho_{RR}(Y)} = \frac{\Pr[A_R(Y) = 1 | \Delta_R(Y) \geq 0, Y]}{\Pr[A_R(Y) = 0 | \Delta_R(Y) \geq 0, Y]}, \quad (27)$$

Suppose first that $Y = S$. In light of our equivalence between prior partisanship and posterior attribution, the denominator of (27) can be understood as the mass of $\mathcal{L}$ voters that switch allegiance and it is maximized when all voters in $crossover_L(Y = S)$ switch. The numerator, in turn is the mass of $\mathcal{R}$ voters that still see R as their favorite and it is minimized when all voters in $crossover_R(Y)$ switch. Thus,

$$\frac{\rho_{RR}(Y = S)}{1 - \rho_{RR}(Y = S)} \geq \frac{\Pr[i \in base_R(Y = S)]}{\Pr[i \in crossover_L(Y = S)]},$$

which gives (2) for $j = R$. A similar derivation would give (2) for $j = L$.

Suppose now that $Y = F$. The denominator of (27) can be understood as the mass of $\mathcal{R}$ voters that keep their allegiance and it is minimized when all voters in $crossover_R(Y = F)$ switch. The numerator, in turn is the fraction of $\mathcal{L}$ voters that switch allegiance after $Y = F$ and it is maximized when all voters in $crossover_L(Y)$ switch. Thus,

$$\frac{\rho_{RR}(Y = F)}{1 - \rho_{RR}(Y = F)} \leq \frac{\Pr[i \in crossover_L(Y = F)]}{\Pr[i \in base_R(Y = F)]}.$$

A similar derivation would apply to $j = L$. Noting that $\rho_{kj}(Y) = 1 - \rho_{jj}(Y)$ gives (3) for $Y = F$. $\square$

Proof of Corollary 1. To simplify notation, let $B_j(Y) \equiv \Pr[i \in base_j(Y)]$ and $C_j(Y) \equiv \Pr[i \in$

$crossover_j(Y)$]. Using the bounds (2) and (3) and the assumption that $C_j(Y) \leq B_k(Y)$ we have

$$\begin{aligned}\rho_{jj}(Y = S) - \rho_{jj}(Y = F) &\geq \frac{1}{1 + \frac{C_k(Y)}{B_j(Y)}} - \frac{1}{1 + \frac{B_j(Y)}{C_k(Y)}} \geq \frac{1}{2} - \frac{1}{2} = 0, \\ \rho_{jj}(Y = S) - \rho_{jk}(Y = S) &\geq \frac{1}{1 + \frac{C_k(Y)}{B_j(Y)}} - \frac{1}{1 + \frac{B_k(Y)}{C_j(Y)}} \geq \frac{1}{2} - \frac{1}{2} = 0, \\ \rho_{jj}(Y = F) - \rho_{jk}(Y = F) &\leq \frac{1}{1 + \frac{B_j(Y)}{C_k(Y)}} - \frac{1}{1 + \frac{C_j(Y)}{B_k(Y)}} \leq \frac{1}{2} - \frac{1}{2} = 0.\end{aligned}$$

□

Proof of Proposition 4. We want to compute the asymptotic belief $\lim_{n \rightarrow \infty} \Pr[a = 1|t(n)]$ when $\lim_{n \rightarrow \infty} t(n) = t$. The proof of Lemma 1 in Acemoglu et al. (2016) (see also Diaconis and Freedman, 1990) shows that after observing in our Bernoulli process outcomes leading to $t(n)$ then $\mathbb{E}^\lambda[f_R^i|t(n)]/\mathbb{E}^\lambda[f_L^i|t(n)] = f_R^i(t)/f_L^i(t)$ so that,

$$\lim_{n \rightarrow \infty} \frac{\Pr_i[a = 1|t(n)]}{\Pr_i[a = 0|t(n)]} = \lim_{n \rightarrow \infty} \frac{\Pr_i[t(n)|a = 1]\alpha_i(1)}{\Pr_i[t(n)|a = 0]\alpha_i(0)} = \frac{f_R^i(t)\alpha_i(1)}{f_L^i(t)\alpha_i(0)}.$$

As priors are stochastically ordered (in the likelihood ratio sense), then $\frac{f_R^i(t)}{f_L^i(t)}$ increases in t for every $i \in \mathcal{R}$, while $\frac{f_R^i(t)}{f_L^i(t)}$ decreases in t for every $i \in \mathcal{L}$. This proves part 1.

For ease of exposition, let $\Lambda_R^i(t) \equiv \frac{f_R^i(t)}{f_L^i(t)}$ and $\Lambda_L^i(t) \equiv \frac{f_L^i(t)}{f_R^i(t)}$. For part 2, take a voter i . After observing $Y_i = Y$, voter i updates attribution according to (11) and (14), with $p_i(Y_i = S) > p_i(Y_i = F)$ if $i \in \mathcal{R}$, but $p_i(Y_i = S) < p_i(Y_i = F)$ if $i \in \mathcal{L}$. If $i \in \mathcal{R}$, then any t such that $\Lambda_R^i(t) \frac{\alpha_i(1)}{\alpha_i(0)} \geq \frac{p_i(Y_i)}{1-p_i(Y_i)}$ will lead the voter to increase their belief that R is actually responsible. If $i \in \mathcal{L}$, then any t such that $\Lambda_L^i(t) \frac{\alpha_i(0)}{\alpha_i(1)} \geq \frac{1-p_i(Y_i)}{p_i(Y_i)}$ will lead the voter to increase their belief that L is responsible. Letting $\tilde{t}_i(Y_i) = (\Lambda_R^i)^{-1}(\frac{\alpha_i(0)}{\alpha_i(1)} \frac{p_i(Y_i)}{1-p_i(Y_i)})$, if $i \in \mathcal{R}$, and $\tilde{t}_i(Y_i) = (\Lambda_L^i)^{-1}(\frac{\alpha_i(1)}{\alpha_i(0)} \frac{1-p_i(Y_i)}{p_i(Y_i)})$, if $i \in \mathcal{L}$, voter i increases her belief that her initially favored politician is responsible iff $t \geq \tilde{t}_i(Y_i)$. □

B Treatment Effects: Bias Structures and Robust Implications

This Appendix discusses systematically how the different treatment-effect tests relate to alternative forms of confounding. The goal is not to introduce new hypotheses, but to clarify which of the theory-implied sign restrictions remain informative under different bias structures.

Recall that the treatment-effect predictions from Proposition 1 in divided regions are:

$$\begin{aligned} HT_1 : \quad & \gamma_4 < 0 \\ HT_2 : \quad & \gamma_4 + \gamma_5 > 0 \\ HT_3 : \quad & \gamma_4 + \gamma_6 > 0 \\ HT_4 : \quad & \gamma_4 + \gamma_5 + \gamma_6 + \gamma_7 < 0. \end{aligned}$$

These correspond directly to the treatment effects in each of the four groups defined by partisanship and assessment.

B.1 Biases in Divided Regions Only

A first concern is that treatment may induce generic shifts in updating that differ across broad respondent categories even in the absence of the mechanism highlighted by the model. Consider the error structure

$$u_i = \zeta_{T\mathcal{R}} + \zeta_{T\mathcal{L}} + \zeta_{TO} + \zeta_{TP} + \epsilon_i, \tag{28}$$

where $\zeta_{T\mathcal{R}}$ and $\zeta_{T\mathcal{L}}$ denote treatment-specific shifts for right-wing and left-wing individuals, while ζ_{TO} and ζ_{TP} denote treatment-specific shifts for good- and bad-assessment individuals.

Under this structure, the primitive tests HT_1 – HT_4 may be contaminated because each group-specific treatment effect combines the theoretical mechanism with these confounder shifts. How-

ever, from Proposition 4 we obtain the following additional restrictions:

$$\begin{aligned}
HE_1 : \quad & \gamma_5 > 0 \\
HE_2 : \quad & \gamma_5 + \gamma_7 < 0 \\
HE_3 : \quad & \gamma_6 > 0 \\
HE_4 : \quad & \gamma_6 + \gamma_7 < 0 \\
HE_5 : \quad & \gamma_7 < 0.
\end{aligned}$$

These restrictions have the advantage of differencing out these confounding shifts:

- HE_1 and HE_2 difference out treatment shifts that are common across ideology groups, but require that treatment-specific shifts by assessment status be symmetric, i.e. $\rho_{TO} - \rho_{TP} = 0$.
- HE_3 and HE_4 difference out treatment shifts that are common across assessment groups, but require that treatment-specific shifts by ideology be symmetric, i.e. $\rho_{TR} - \rho_{TL} = 0$.
- HE_5 is the most robust implication in this setup, since it differences out both types of additive treatment-specific shifts and does not require either symmetry condition.

Thus, in the within-divided-regions specification, the primitive hypotheses are the most direct tests of Proposition 1, but the additional implications are useful because they remain informative under broader classes of confounding.

B.2 Biases Common to Divided and United Regions

A second concern is that treatment responsiveness may differ flexibly across the four respondent categories, but in a way that is common across divided and united regions. Consider the error structure

$$u_i = \zeta_{TLP} + \zeta_{TRP} + \zeta_{TLO} + \zeta_{TRO} + \epsilon_i, \quad (29)$$

where, for example, ζ_{TLP} denotes a treatment-specific shift for left-wing bad-assessment individuals unrelated to the model.

Under this structure, the fourth-difference specification,

$$\begin{aligned}
P_i = & \delta_0 \\
& + \delta_1 O_i + \delta_2 \mathcal{R}_i + \delta_3 O_i \times \mathcal{R}_i \\
& + \delta_4 T_i + \delta_5 O_i \times T_i + \delta_6 \mathcal{R}_i \times T_i + \delta_7 O_i \times \mathcal{R}_i \times T_i \\
& + \delta_8 D_{ig} \\
& + \delta_9 O_i \times D_{ig} + \delta_{10} \mathcal{R}_i \times D_{ig} + \delta_{11} O_i \times \mathcal{R}_i \times D_{ig} \\
& + \delta_{12} T_i \times D_{ig} + \delta_{13} O_i \times T_i \times D_{ig} + \delta_{14} \mathcal{R}_i \times T_i \times D_{ig} + \delta_{15} O_i \times \mathcal{R}_i \times T_i \times D_{ig} \\
& + \phi_g + u_i,
\end{aligned}$$

differences out those common category-specific treatment shifts. The primitive restrictions are then:

$$\begin{aligned}
HT_1 : & \quad \delta_{12} < 0 \\
HT_2 : & \quad \delta_{12} + \delta_{13} > 0 \\
HT_3 : & \quad \delta_{12} + \delta_{14} > 0 \\
HT_4 : & \quad \delta_{12} + \delta_{13} + \delta_{14} + \delta_{15} < 0.
\end{aligned}$$

Under the identifying assumption that the idiosyncratic updating by category is common across divided and united regions, these primitive fourth-difference tests are already robust to the confounding structure in (29).

B.3 Biases that Differ Across Region Types

A third concern is that some broad treatment-specific shifts may themselves differ between divided and united regions. Consider the error structure

$$u_i = \zeta_{D\mathcal{RT}} + \zeta_{D\mathcal{LT}} + \zeta_{DOT} + \zeta_{DPT} + \epsilon_i, \quad (30)$$

where $\zeta_{D\mathcal{RT}}$ captures, for example, a treatment-specific shift for right-wing individuals in divided regions relative to right-wing individuals in united regions, unrelated to the model.

In this case, the primitive fourth-difference restrictions may still be contaminated. However, the same sign pattern implies the following auxiliary restrictions:

$$\begin{aligned}
HE_1 : \quad & \delta_{13} > 0 \\
HE_2 : \quad & \delta_{13} + \delta_{15} < 0 \\
HE_3 : \quad & \delta_{14} > 0 \\
HE_4 : \quad & \delta_{14} + \delta_{15} < 0 \\
HE_5 : \quad & \delta_{15} < 0.
\end{aligned}$$

Again, these restrictions differ in robustness:

- HE_1 and HE_2 remain valid if treatment-specific shifts by assessment status are symmetric across divided and united regions, i.e. $\rho_{DOT} - \rho_{DPT} = 0$.
- HE_3 and HE_4 remain valid if treatment-specific shifts by ideology are symmetric across divided and united regions, i.e. $\rho_{D\mathcal{RT}} - \rho_{D\mathcal{LT}} = 0$.
- HE_5 is again the strongest implication, since it differences out both types of additive shifts and does not require either symmetry condition.

B.4 Interpretation

The discussion above clarifies the role of the different treatment-effect tests. The primitive hypotheses HT_1 – HT_4 are the most direct empirical translation of Proposition 1. The auxiliary implications HE_1 – HE_5 follow from Proposition 4 and are useful because some of them remain informative under richer classes of confounding. In particular, the highest-order interaction restrictions, γ_7 in the divided-regions specification and δ_{15} in the fourth-difference specification, provide the most robust falsification test of the theory’s sign pattern.

C Appendix Tables and Figures

Table C.1**Government Coalitions and Divided Governments, by Region**

Region	President	Gov Coalition	Gov Formation	Divided Gov
	(1)	(2)	(3)	(4)
Central Government	PSOE	PSOE, UP	PSOE, UP, MP, PNV, BNG, Reg.	
Andalusia	PP	PP, Cs	PP, Cs, VOX	Yes
Aragon	PSOE	PSOE, UP, Reg.	PSOE, UP, Reg.	No
Asturias	PSOE	PSOE	PSOE, UP	No
Canary Islands	PSOE	PSOE, UP, Reg.	PSOE, UP, Reg.	No
Cantabria	Reg	PSOE, Reg.	PSOE, Reg.	No
Castile and Leon	PP	PP, Cs	PP, Cs	Yes
Castile-La Mancha	PSOE	PSOE	PSOE	No
Catalonia	ERC	JxC, ERC	JxC, ERC	Yes
Valencian Com.	PSOE	PSOE, UP, Reg.	PSOE, UP, Reg.	No
Com. Madrid	PP	PP, Cs	PP, Cs, VOX	Yes
Galicia	PP	PP	PP	Yes
Extremadura	PSOE	PSOE	PSOE	No
Balearic Islands	PSOE	PSOE, UP, Reg.	PSOE, UP, Reg.	No
La Rioja	PSOE	PSOE, UP	PSOE, UP	No
Murcia	PP	PP, Cs	PP, Cs, VOX	Yes
Navarre	PSOE	PSOE, UP, PNV	PSOE, UP, PNV	No
Basque Country	PNV	PNV, PSOE	PNV, PSOE	No

Notes: The column President indicates the party of the president of the regional government at the time of the survey. Gov Coalition indicates the parties that are part of the regional government (the executive). Gov Formation indicates the parties that voted “yes” to the investiture of the regional president. The first row shows analogous values for the central government, i.e., party of the prime minister, parties that are part of the central government, and parties that voted “yes” to the investiture of the prime minister. Divided Gov indicates regions in which there is no overlap between the government formation parties for that region and for the central government. “Reg.” stands for voting for any regionalist (or nationalist) party running in that region only. Our sample does not contain any respondent from the autonomous cities of Ceuta and Melilla.

Table C.2**Sample Characteristics**

	Spanish Population	
	(source: INE)	Our Sample
Female	0.52	0.56
Ages 18-24	0.08	0.07
Ages: 25-34	0.14	0.14
Ages: 35-44	0.19	0.20
Ages: 45-54	0.19	0.21
Ages: 55+	0.39	0.38
North-East Region	0.21	0.21
East Region	0.14	0.14
South Region	0.24	0.23
Center Region	0.22	0.24
North-West Region	0.09	0.09
North Region	0.09	0.08
Primary Education or Less	0.18	0.02
Secondary Education	0.29	0.11
Upper Secondary Education	0.14	0.23
Vocational Training	0.08	0.30
Tertiary Education	0.31	0.33

Notes: Demographic characteristics of the survey sample compared with the Spanish adult population. Population statistics are from the National Institute of Statistics (Instituto Nacional de Estadística, INE), 2019. Survey data are from the authors' sample of 3,857 respondents collected July 10–17, 2023. The survey was conducted online by YouGov using quota sampling to achieve representativeness by age, gender, region, and education. North-East Region includes Catalonia, Aragon, and Balearic Islands; East Region includes Valencia and Murcia; South Region includes Andalusia and Canary Islands; Center Region includes Madrid, Castile-La Mancha, Castile and León, and Extremadura; North-West Region includes Galicia, Asturias, and Cantabria; North Region includes Basque Country, Navarra, and La Rioja.

Table C.3

Summary Statistics

	Mean	Min.	Max.	Std. Dev.	Observations
Age <35	0.22	0	1	0.42	3857
Age 36-50	0.32	0	1	0.47	3857
Age >50	0.45	0	1	0.50	3857
Low Educ	0.14	0	1	0.35	3857
High Educ	0.86	0	1	0.35	3857
Female	0.56	0	1	0.50	3857
Kids	0.34	0	1	0.47	3857
Rural	0.22	0	1	0.42	3857
Worker	0.61	0	1	0.49	3857
Student	0.06	0	1	0.24	3857
Retired	0.14	0	1	0.35	3857
Good Assessment	0.39	0	1	0.49	3857
Good Assessment (Continuous)	-40.00	-910	122	103.04	3857
Right-wing Voter	0.46	0	1	0.50	3806
Divided Region	0.65	0	1	0.48	3857
Treated	0.67	0	1	0.47	3857
Resp RG (vs CG)	1.37	-10	10	6.19	3857
DeJure RG	0.58	0	1	0.49	3857

Notes: Summary statistics for the main variables used in the analysis. The unit of observation is an individual survey respondent. “Good Assessment” indicates respondents whose prior beliefs about waiting times are below the official regional average. “Good Assessment (Continuous)” is the difference between official and believed waiting times in days. “Right Voter” indicates right-wing partisanship based on past regional vote or ideological self-placement. “Divided Region” indicates residence in a region where the regional government coalition does not overlap with the central government coalition. “Treated” indicates assignment to receive accurate information about regional waiting times. “Responsibility Regional Government (vs. Central Government)” is measured on a scale from -10 (central government fully responsible) to $+10$ (regional government fully responsible).

Table C.4

Attribution of Responsibility: Observational Evidence

	(1)	(2)	(3)
	Divided Reg.	United Reg.	All
Good Assessment	-1.09** (0.59)	0.98 (0.81)	0.98 (0.80)
Right-wing Voter	-3.83*** (0.58)	-1.77*** (0.73)	-1.77*** (0.73)
Good Assessment $\times$ Right-wing Voter	2.78*** (0.89)	-0.60 (1.18)	-0.60 (1.17)
Good Assessment $\times$ Divided Region			-2.07** (1.00)
Right-wing Voter $\times$ Divided Region			-2.06** (0.93)
Good Assessment $\times$ Right-wing Voter $\times$ Divided Region			3.37** (1.47)
Constant	2.60*** (0.37)	2.54*** (0.54)	2.58*** (0.30)
H1<0	-1.09 [0.03]	0.98 [0.89]	-2.07 [0.02]
H2>0	1.69 [0.01]	0.39 [0.33]	1.30 [0.12]
H3<0	-3.83 [0.00]	-1.77 [0.01]	-2.06 [0.01]
H4>0	-1.05 [0.94]	-2.36 [0.99]	1.31 [0.12]
H5>0	2.78 [0.00]	-0.60 [0.69]	3.37 [0.01]
Observations	819	434	1253
R^2	0.06	0.06	0.07

Notes: Sample restricted to respondents in the control group. Columns (1) and (2) show the results from estimating equation (4) for respondents in divided and united regions, respectively. Column (3) shows the results from estimating equation (5) for the full sample. A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of divided regions. The dependent variable is how much responsibility the respondent assigns to the regional versus the central government, measured on a scale from -10 (all responsibility lies with the central government) to 10 (all responsibility lies with the regional government). “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Bad Assessment” and “Good Assessment” indicate whether respondents’ prior beliefs about waiting times exceed or fall below the official regional average, respectively. Rows below the constant show results from testing the hypotheses outlined in Section 2, one-sided p-values of the tests are in square brackets. Robust standard errors in parentheses. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table C.5

De Jure Responsibility: Observational Evidence

	(1)	(2)	(3)
	Divided Reg.	United Reg.	All
Good Assessment	0.01 (0.04)	-0.08 (0.07)	-0.08 (0.07)
Right-wing Voter	-0.16*** (0.05)	-0.16*** (0.06)	-0.16*** (0.06)
Good Assessment $\times$ Right-wing Voter	0.08 (0.07)	0.02 (0.10)	0.02 (0.10)
Good Assessment $\times$ Divided Region			0.09 (0.08)
Right-wing Voter $\times$ Divided Region			0.00 (0.07)
Good Assessment $\times$ Right-wing Voter $\times$ Divided Region			0.06 (0.12)
Constant	0.64*** (0.03)	0.69*** (0.04)	0.65*** (0.02)
H1<0	0.01 [0.55]	-0.08 [0.13]	0.09 [0.85]
H2>0	0.08 [0.07]	-0.06 [0.80]	0.15 [0.06]
H3<0	-0.16 [0.00]	-0.16 [0.00]	0.00 [0.51]
H4>0	-0.08 [0.93]	-0.14 [0.95]	0.06 [0.27]
H5>0	0.08 [0.14]	0.02 [0.43]	0.06 [0.32]
Observations	819	434	1253
R^2	0.03	0.06	0.04

Notes: Sample restricted to respondents in the control group. Columns (1) and (2) show the results from estimating equation (4) for respondents in divided and united regions, respectively. Column (3) shows the results from estimating equation (5) for the full sample. A region is classified as "divided" if none of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of divided regions. The dependent variable is a dummy variable indicating if the respondent (correctly) identified that regions are *legally* responsible for managing healthcare services in Spain. "Left-wing" and "Right-wing" refer to respondents' partisan affiliation based on past regional vote or ideological self-placement. "Bad Assessment" and "Good Assessment" indicate whether respondents' prior beliefs about waiting times exceed or fall below the official regional average, respectively. Rows below the constant show results from testing the hypotheses outlined in Section 2, one-sided p-values of the tests are in square brackets. Robust standard errors in parentheses. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table C.6

Motivated Learning in United Regions

	(1)	(2)	(3)	(4)
	Flexible, treated	Flexible, treated vs control	Parsimonious, treated	Parsimonious, treated vs control
Right-wing voter	7.14 (7.84)	-8.80 (13.19)	4.96 (7.68)	-10.05 (12.90)
Bad-news gap (Official – Prior)	0.75*** (0.07)	0.43*** (0.09)	0.77*** (0.07)	0.44*** (0.09)
Right-wing voter × Bad-news gap	0.03 (0.10)	0.01 (0.13)		
Bad-news indicator	22.46* (14.12)	-1.53 (15.20)	12.61* (8.24)	-7.07 (12.49)
Right-wing voter=1 × Bad-news indicator=1	-29.23** (17.60)	-18.97 (19.13)	-6.14 (9.63)	-6.00 (13.32)
Bad-news gap × Bad-news indicator	-0.22 (0.28)	-0.13 (0.25)		
Right-wing voter × Bad-news gap × Bad-news indicator	0.53* (0.37)	0.32 (0.33)		
Treatment		-16.98* (11.69)		-16.22* (11.40)
Treatment=1 × Right-wing voter=1		17.06 (15.39)		16.11 (15.05)
Treatment × Bad-news gap		0.32*** (0.12)		0.34*** (0.11)
Right-wing voter × Treatment × Bad-news gap		0.03 (0.16)		
Treatment=1 × Bad-news indicator=1		22.65 (20.55)		18.20 (14.98)
Treatment=1 × Right-wing voter=1 × Bad-news indicator=1		-10.13 (25.73)		0.25 (16.43)
Treatment × Bad-news gap × Bad-news indicator		-0.09 (0.37)		
Right-wing voter × Treatment × Bad-news gap × Bad-news indicator		0.21 (0.49)		
Congruent news × Bad-news gap			-0.00 (0.10)	0.01 (0.12)
Treatment × Congruent news × Bad-news gap				-0.02 (0.16)
Observations	887	1321	887	1321
R^2	0.56	0.48	0.55	0.48

Notes: The dependent variable is the difference between posterior and prior beliefs about waiting times (in days). “Bad-news gap” is defined as official waiting times minus prior beliefs (in days). “Bad-news indicator” takes the value of 1 if official waiting time is larger than prior waiting time. Columns 1 and 2 report flexible specifications that allow the learning slope to differ across right- and left-wing respondents and across bad- and good-news signals. Column 1 is estimated only among treated respondents, while column 2 includes both treated and control respondents and identifies the asymmetric updating pattern relative to the control group. Columns 3 and 4 report specifications based on a congruent-news indicator, equal to one for right-wing respondents receiving bad news and left-wing respondents receiving good news. Column 3 is estimated only among treated respondents, while column 4 estimates the corresponding treatment-versus-control specification. A region is classified as “united” if any of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for the classification of regions. “Right-wing” refers to respondents’ partisan affiliation based on past regional vote or ideological self-placement. All specifications include region fixed effects. Robust standard errors in parentheses. Significance levels for one-sided hypothesis testing: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table C.7

Attribution of Responsibility: Experimental Evidence

	(1)	(2)	(3)
	Divided Reg.	United Reg.	All
Good Assessment	-1.09** (0.58)	1.01 (0.79)	1.01* (0.78)
Right-wing Voter	-3.80*** (0.57)	-1.65** (0.73)	-1.65** (0.73)
Good Assessment × Right-wing Voter	2.80*** (0.89)	-0.67 (1.15)	-0.67 (1.15)
Treated	0.45 (0.43)	0.43 (0.63)	0.43 (0.63)
Good Assessment × Treated	0.01 (0.70)	-2.35*** (0.98)	-2.35*** (0.98)
Right-wing Voter × Treated	0.31 (0.70)	-0.24 (0.88)	-0.24 (0.88)
Good Assessment × Right-wing Voter × Treated	-1.88** (1.08)	1.37 (1.41)	1.37 (1.41)
Good Assessment × Divided Region			-2.10** (0.98)
Right-wing Voter × Divided Region			-2.15** (0.93)
Good Assessment × Right-wing Voter × Divided Region			3.47*** (1.45)
Treated × Divided Region			0.01 (0.76)
Good Assessment × Treated × Divided Region			2.35** (1.20)
Right-wing Voter × Treated × Divided Region			0.55 (1.12)
Good Assessment × Right-wing Voter × Treated × Divided Region			-3.25** (1.78)
Constant	2.58*** (0.36)	2.46*** (0.53)	2.54*** (0.30)
HT1<0	0.45 [0.85]	0.43 [0.76]	0.01 [0.51]
HT2>0	0.45 [0.20]	-1.91 [0.99]	2.37 [0.01]
HT3>0	0.76 [0.08]	0.19 [0.38]	0.56 [0.25]
HT4<0	-1.12 [0.04]	-0.79 [0.17]	-0.33 [0.37]
HE1>0	0.01 [0.50]	-2.35 [0.99]	2.35 [0.03]
HE2<0	-1.87 [0.01]	-0.98 [0.17]	-0.90 [0.25]
HE3>0	0.31 [0.33]	-0.24 [0.61]	0.55 [0.31]
HE4<0	-1.57 [0.03]	1.13 [0.85]	-2.70 [0.03]
HE5<0	-1.88 [0.04]	1.37 [0.83]	-3.25 [0.03]
Observations	2485	1321	3806
R ²	0.07	0.05	0.06

Notes: Columns (1) and (2) show the results from estimating equation (6) for respondents in divided and united regions, respectively. Column (3) shows the results from estimating equation (8) for the full sample. A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of divided regions. The dependent variable is how much responsibility the respondent assigns to the regional versus the central government, measured on a scale from -10 (all responsibility lies with the central government) to 10 (all responsibility lies with the regional government). “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Bad Assessment” and “Good Assessment” indicate whether respondents’ prior beliefs about waiting times exceed or fall below the official regional average, respectively. Rows below the constant show results from testing the hypotheses outlined in Section 2, one-sided p-values of the tests are in square brackets. Robust standard errors in parentheses. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table C.8

De Jure Responsibility: Experimental Evidence

	(1)	(2)	(3)
	Divided Reg.	United Reg.	All
Good Assessment	0.00 (0.04)	-0.09 (0.07)	-0.09 (0.07)
Right-wing Voter	-0.16*** (0.05)	-0.16*** (0.06)	-0.16*** (0.06)
Good Assessment × Right-wing Voter	0.07 (0.07)	0.00 (0.10)	0.00 (0.10)
Treated	0.03 (0.03)	-0.03 (0.05)	-0.03 (0.05)
Good Assessment × Treated	-0.06 (0.05)	0.07 (0.09)	0.07 (0.09)
Right-wing Voter × Treated	-0.02 (0.06)	0.03 (0.07)	0.03 (0.07)
Good Assessment × Right-wing Voter × Treated	0.01 (0.09)	0.04 (0.12)	0.04 (0.12)
Good Assessment × Divided Region			0.09 (0.08)
Right-wing Voter × Divided Region			0.00 (0.07)
Good Assessment × Right-wing Voter × Divided Region			0.07 (0.12)
Treated × Divided Region			0.06 (0.06)
Good Assessment × Treated × Divided Region			-0.12 (0.10)
Right-wing Voter × Treated × Divided Region			-0.05 (0.09)
Good Assessment × Right-wing Voter × Treated × Divided Region			-0.03 (0.15)
Constant	0.64*** (0.03)	0.69*** (0.04)	0.66*** (0.02)
HT1<0	0.03 [0.82]	-0.03 [0.27]	0.06 [0.85]
HT2>0	-0.03 [0.73]	0.04 [0.31]	-0.06 [0.77]
HT3>0	0.01 [0.42]	0.00 [0.49]	0.01 [0.45]
HT4<0	-0.04 [0.23]	0.10 [0.93]	-0.14 [0.05]
HE1>0	-0.06 [0.85]	0.07 [0.22]	-0.12 [0.89]
HE2<0	-0.05 [0.24]	0.10 [0.89]	-0.15 [0.08]
HE3>0	-0.02 [0.65]	0.03 [0.32]	-0.05 [0.73]
HE4<0	-0.01 [0.43]	0.07 [0.75]	-0.08 [0.25]
HE5<0	0.01 [0.54]	0.04 [0.61]	-0.03 [0.43]
Observations	2485	1321	3806
R ²	0.03	0.04	0.03

Notes: Columns (1) and (2) show the results from estimating equation (6) for respondents in divided and united regions, respectively. Column (3) shows the results from estimating equation (8) for the full sample. A region is classified as "divided" if none of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of divided regions. The dependent variable is a dummy variable indicating if the respondent (correctly) identified that regions are *legally* responsible for managing healthcare services in Spain. "Left-wing" and "Right-wing" refer to respondents' partisan affiliation based on past regional vote or ideological self-placement. "Bad Assessment" and "Good Assessment" indicate whether respondents' prior beliefs about waiting times exceed or fall below the official regional average, respectively. Rows below the constant show results from testing the hypotheses outlined in Section 2, one-sided p-values of the tests are in square brackets. Robust standard errors in parentheses. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table C.9

Divided Regions, Observational Evidence: Balance

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)
	Age <35	Age 36-50	Age >50	Low Educ	High Educ	Female	Kids	Rural	Worker	Student	Retired
Good Assessment	0.02 (0.04)	-0.06 (0.04)	0.04 (0.05)	0.00 (0.03)	-0.00 (0.03)	-0.08* (0.05)	0.01 (0.05)	-0.04 (0.04)	0.07 (0.05)	0.02 (0.02)	0.01 (0.03)
Right-wing Voter	-0.00 (0.04)	-0.04 (0.04)	0.05 (0.05)	-0.01 (0.03)	0.01 (0.03)	0.03 (0.05)	-0.02 (0.04)	0.06* (0.04)	0.04 (0.04)	0.00 (0.02)	-0.01 (0.03)
Good Assessment × Right-wing Voter	-0.03 (0.06)	0.10 (0.07)	-0.07 (0.07)	-0.03 (0.05)	0.03 (0.05)	-0.08 (0.07)	0.01 (0.07)	0.03 (0.06)	-0.14** (0.07)	-0.01 (0.04)	0.03 (0.05)
Constant	0.23*** (0.02)	0.34*** (0.03)	0.43*** (0.03)	0.14*** (0.02)	0.86*** (0.02)	0.59*** (0.03)	0.34*** (0.03)	0.19*** (0.02)	0.59*** (0.03)	0.06*** (0.01)	0.15*** (0.02)
Observations	819	819	819	819	819	819	819	819	819	819	819
R^2	0.00	0.01	0.01	0.01	0.01	0.03	0.01	0.10	0.01	0.01	0.01

Notes: Results from estimating equation (4). Sample restricted to respondents in divided regions and the control group. Each column uses a demographic characteristic as the dependent variable: age groups (columns 1–3), education levels (columns 4–5), female (column 6), having children (column 7), rural residence (column 8), and employment status (columns 9–11). A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of divided regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Bad Assessment” and “Good Assessment” indicate whether respondents’ prior beliefs about waiting times exceed or fall below the official regional average, respectively. Robust standard errors in parentheses. Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table C.10

Divided Regions Controlling for United, Observational Evidence: Balance

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)
	Age <35	Age 36-50	Age >50	Low Educ	High Educ	Female	Kids	Rural	Worker	Student	Retired
Good Assessment	-0.03 (0.06)	0.06 (0.07)	-0.02 (0.08)	0.02 (0.05)	-0.02 (0.05)	-0.17** (0.07)	0.08 (0.07)	0.02 (0.06)	-0.09 (0.08)	-0.03 (0.04)	-0.02 (0.06)
Right-wing Voter	0.02 (0.05)	0.01 (0.06)	-0.03 (0.06)	-0.01 (0.04)	0.01 (0.04)	-0.08 (0.06)	0.03 (0.06)	0.01 (0.05)	-0.01 (0.06)	-0.00 (0.03)	-0.04 (0.04)
Good Assessment × Right-wing Voter	0.15* (0.09)	-0.08 (0.10)	-0.07 (0.10)	0.00 (0.08)	-0.00 (0.08)	0.07 (0.10)	-0.02 (0.10)	0.02 (0.09)	0.15 (0.10)	0.03 (0.06)	-0.02 (0.07)
Good Assessment × Divided Region	0.06 (0.07)	-0.12 (0.08)	0.06 (0.09)	-0.02 (0.06)	0.02 (0.06)	0.09 (0.09)	-0.07 (0.08)	-0.06 (0.07)	0.15* (0.09)	0.05 (0.05)	0.04 (0.07)
Right-wing Voter × Divided Region	-0.02 (0.06)	-0.05 (0.07)	0.07 (0.08)	0.00 (0.05)	-0.00 (0.05)	0.11 (0.07)	-0.05 (0.07)	0.05 (0.06)	0.06 (0.07)	0.00 (0.04)	0.04 (0.05)
Good Assessment × Right-wing Voter × Divided Region	-0.18* (0.11)	0.18 (0.12)	0.00 (0.13)	-0.03 (0.09)	0.03 (0.09)	-0.16 (0.12)	0.03 (0.12)	0.01 (0.10)	-0.29** (0.13)	-0.05 (0.07)	0.05 (0.08)
Constant	0.22*** (0.02)	0.33*** (0.02)	0.45*** (0.02)	0.14*** (0.02)	0.86*** (0.02)	0.62*** (0.02)	0.33*** (0.02)	0.21*** (0.02)	0.60*** (0.02)	0.07*** (0.01)	0.15*** (0.02)
Observations	1,253	1,253	1,253	1,253	1,253	1,253	1,253	1,253	1,253	1,253	1,253
R^2	0.01	0.01	0.01	0.02	0.02	0.04	0.02	0.10	0.01	0.02	0.02

Notes: Results from estimating equation (5). Sample restricted to respondents in the control group. Each column uses a demographic characteristic as the dependent variable: age groups (columns 1–3), education levels (columns 4–5), female (column 6), having children (column 7), rural residence (column 8), and employment status (columns 9–11). A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of divided regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Bad Assessment” and “Good Assessment” indicate whether respondents’ prior beliefs about waiting times exceed or fall below the official regional average, respectively. Robust standard errors in parentheses. Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table C.11

Divided Regions, Experimental Results: Balance

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)
	Age <35	Age 36-50	Age >50	Low Educ	High Educ	Female	Kids	Rural	Worker	Student	Retired
Good Assessment	0.02 (0.04)	-0.06 (0.04)	0.04 (0.05)	0.00 (0.03)	-0.00 (0.03)	-0.08* (0.05)	0.00 (0.04)	-0.04 (0.04)	0.07 (0.05)	0.03 (0.02)	0.01 (0.03)
Right-wing Voter	-0.01 (0.04)	-0.04 (0.04)	0.05 (0.05)	-0.00 (0.03)	0.00 (0.03)	0.03 (0.04)	-0.01 (0.04)	0.06 (0.03)	0.05 (0.04)	0.00 (0.02)	-0.01 (0.03)
Good Assessment $\times$ Right-wing Voter	-0.04 (0.06)	0.10 (0.07)	-0.07 (0.07)	-0.02 (0.05)	0.02 (0.05)	-0.09 (0.07)	0.02 (0.07)	0.03 (0.06)	-0.14* (0.07)	-0.01 (0.04)	0.03 (0.05)
Treated	0.00 (0.03)	-0.03 (0.03)	0.03 (0.04)	-0.02 (0.02)	0.02 (0.02)	-0.01 (0.03)	-0.02 (0.03)	0.02 (0.03)	0.02 (0.03)	0.01 (0.02)	-0.01 (0.03)
Good Assessment $\times$ Treated	0.01 (0.05)	0.04 (0.05)	-0.05 (0.06)	0.03 (0.04)	-0.03 (0.04)	-0.02 (0.06)	0.01 (0.05)	0.04 (0.04)	-0.04 (0.06)	-0.02 (0.03)	-0.02 (0.04)
Right-wing Voter $\times$ Treated	-0.02 (0.05)	0.06 (0.05)	-0.04 (0.06)	0.01 (0.04)	-0.01 (0.04)	0.03 (0.06)	0.06 (0.05)	-0.05 (0.04)	-0.04 (0.06)	-0.03 (0.03)	-0.00 (0.04)
Good Assessment $\times$ Right-wing Voter $\times$ Treated	0.04 (0.07)	-0.07 (0.08)	0.04 (0.09)	-0.00 (0.06)	0.00 (0.06)	0.03 (0.09)	-0.06 (0.08)	-0.03 (0.07)	0.10 (0.09)	0.03 (0.04)	0.00 (0.06)
Constant	0.23*** (0.02)	0.34*** (0.03)	0.43*** (0.03)	0.14*** (0.02)	0.86*** (0.02)	0.59*** (0.03)	0.34*** (0.03)	0.19*** (0.02)	0.58*** (0.03)	0.06*** (0.01)	0.15*** (0.02)
Observations	2,485	2,485	2,485	2,485	2,485	2,485	2,485	2,485	2,485	2,485	2,485
R^2	0.01	0.01	0.01	0.01	0.01	0.02	0.01	0.08	0.00	0.01	0.01

Notes: Results from estimating equation (6). Sample restricted to respondents in divided regions. Each column uses a demographic characteristic as the dependent variable: age groups (columns 1–3), education levels (columns 4–5), female (column 6), having children (column 7), rural residence (column 8), and employment status (columns 9–11). A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of divided regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Bad Assessment” and “Good Assessment” indicate whether respondents’ prior beliefs about waiting times exceed or fall below the official regional average, respectively. Robust standard errors in parentheses. Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table C.12

Divided Regions Controlling for United, Experimental Results: Balance

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)
	Age <35	Age 36-50	Age >50	Low Educ	High Educ	Female	Kids	Rural	Worker	Student	Retired
Good Assessment	-0.03 (0.06)	0.06 (0.07)	-0.03 (0.08)	0.03 (0.05)	-0.03 (0.05)	-0.18** (0.07)	0.09 (0.07)	0.02 (0.06)	-0.10 (0.07)	-0.02 (0.04)	-0.03 (0.05)
Right-wing Voter	0.02 (0.05)	0.01 (0.06)	-0.03 (0.06)	-0.01 (0.04)	0.01 (0.04)	-0.09 (0.06)	0.04 (0.06)	0.01 (0.05)	-0.02 (0.06)	0.01 (0.03)	-0.04 (0.04)
Good Assessment × Right-wing Voter	0.14 (0.09)	-0.07 (0.10)	-0.07 (0.10)	0.01 (0.07)	-0.01 (0.07)	0.09 (0.10)	-0.03 (0.10)	0.01 (0.09)	0.15 (0.10)	0.03 (0.06)	-0.02 (0.07)
Good Assessment × Divided Region	0.05 (0.07)	-0.12 (0.08)	0.07 (0.09)	-0.03 (0.06)	0.03 (0.06)	0.10 (0.09)	-0.09 (0.08)	-0.06 (0.07)	0.17* (0.09)	0.05 (0.04)	0.04 (0.06)
Right-wing Voter × Divided Region	-0.03 (0.06)	-0.05 (0.07)	0.08 (0.07)	0.00 (0.05)	-0.00 (0.05)	0.12 (0.07)	-0.05 (0.07)	0.05 (0.06)	0.06 (0.07)	-0.00 (0.04)	0.03 (0.05)
Good Assessment × Right-wing Voter × Divided Region	-0.17 (0.11)	0.17 (0.12)	0.00 (0.12)	-0.03 (0.09)	0.03 (0.09)	-0.17 (0.12)	0.05 (0.12)	0.03 (0.10)	-0.29** (0.12)	-0.04 (0.07)	0.05 (0.08)
Treated	0.00 (0.04)	0.02 (0.05)	-0.02 (0.06)	-0.00 (0.04)	0.00 (0.04)	-0.01 (0.05)	0.01 (0.05)	0.01 (0.04)	0.04 (0.05)	-0.06** (0.03)	-0.03 (0.04)
Good Assessment × Treated	0.07 (0.07)	-0.10 (0.09)	0.03 (0.09)	-0.04 (0.06)	0.04 (0.06)	-0.03 (0.09)	-0.08 (0.09)	-0.04 (0.07)	-0.09 (0.09)	0.10** (0.05)	0.07 (0.07)
Right-wing Voter × Treated	-0.03 (0.06)	-0.02 (0.07)	0.05 (0.07)	0.02 (0.05)	-0.02 (0.05)	-0.02 (0.07)	0.02 (0.07)	0.00 (0.06)	-0.04 (0.07)	0.02 (0.03)	0.06 (0.05)
Good Assessment × Right-wing Voter × Treated	-0.15 (0.10)	0.14 (0.12)	0.01 (0.12)	0.05 (0.09)	-0.05 (0.09)	0.03 (0.12)	0.05 (0.12)	0.06 (0.10)	0.03 (0.12)	-0.07 (0.07)	-0.02 (0.08)
Divided Region × Treated	-0.00 (0.05)	-0.04 (0.06)	0.05 (0.07)	-0.02 (0.05)	0.02 (0.05)	0.01 (0.06)	-0.04 (0.06)	0.01 (0.05)	-0.02 (0.06)	0.07** (0.03)	0.02 (0.05)
Good Assessment × Divided Region × Treated	-0.06 (0.09)	0.15 (0.10)	-0.08 (0.11)	0.06 (0.08)	-0.06 (0.08)	0.01 (0.11)	0.10 (0.10)	0.08 (0.09)	0.05 (0.11)	-0.13** (0.05)	-0.09 (0.08)
Right-wing Voter × Divided Region × Treated	0.01 (0.07)	0.08 (0.09)	-0.09 (0.09)	-0.01 (0.06)	0.01 (0.06)	0.05 (0.09)	0.04 (0.09)	-0.05 (0.07)	-0.00 (0.09)	-0.05 (0.04)	-0.06 (0.06)
Good Assessment × Right-wing Voter × Divided Region × Treated	0.19 (0.13)	-0.22 (0.14)	0.03 (0.15)	-0.05 (0.11)	0.05 (0.11)	-0.00 (0.15)	-0.11 (0.15)	-0.09 (0.13)	0.07 (0.15)	0.10 (0.08)	0.02 (0.10)
Constant	0.22*** (0.02)	0.33*** (0.02)	0.45*** (0.02)	0.14*** (0.02)	0.86*** (0.02)	0.62*** (0.02)	0.32*** (0.02)	0.21*** (0.02)	0.60*** (0.02)	0.06*** (0.01)	0.15*** (0.02)
Observations	3,806	3,806	3,806	3,806	3,806	3,806	3,806	3,806	3,806	3,806	3,806
R^2	0.01	0.01	0.01	0.01	0.01	0.02	0.01	0.09	0.01	0.01	0.01

Notes: Results from estimating equation (8). Each column uses a demographic characteristic as the dependent variable: age groups (columns 1–3), education levels (columns 4–5), female (column 6), having children (column 7), rural residence (column 8), and employment status (columns 9–11). A region is classified as “divided” if none of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of divided regions. “Left-wing” and “Right-wing” refer to respondents’ partisan affiliation based on past regional vote or ideological self-placement. “Bad Assessment” and “Good Assessment” indicate whether respondents’ prior beliefs about waiting times exceed or fall below the official regional average, respectively. Robust standard errors in parentheses. Significance levels: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

Table C.13

Divided Regions, Observational Evidence: Robustness

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)
	Base.	No FE	Controls	Right: ideol.	Right: vote	Opt.: median	Opt.: ideal	Opt.: reg.-ideol.	Opt.: cont.	No Catalonia
Good Assessment	-1.09** (0.59)	-1.23** (0.59)	-1.13** (0.58)	-0.44 (0.65)	-1.07* (0.67)	-1.64*** (0.57)	-1.66*** (0.57)	-1.52*** (0.57)	-0.39 (0.30)	-0.95* (0.73)
Right-wing Voter	-3.83*** (0.58)	-3.66*** (0.57)	-3.94*** (0.57)	-3.63*** (0.56)	-4.63*** (0.76)	-4.07*** (0.63)	-3.84*** (0.64)	-4.03*** (0.61)	-2.82*** (0.45)	-4.06*** (0.65)
Good Assessment $\times$ Right-wing Voter	2.78*** (0.89)	2.88*** (0.89)	2.78*** (0.88)	1.13* (0.88)	2.55** (1.12)	2.78*** (0.88)	2.43*** (0.88)	2.80*** (0.87)	1.39*** (0.50)	2.17** (1.02)
Constant	2.60*** (0.37)	2.56*** (0.37)	2.66*** (0.36)	3.01*** (0.39)	2.79*** (0.42)	2.94*** (0.40)	2.88*** (0.38)	2.91*** (0.40)	2.16*** (0.29)	3.00*** (0.44)
H1<0	-1.09**	-1.23**	-1.13**	-0.44	-1.07*	-1.64***	-1.66***	-1.52***	-0.39	-0.95*
H2>0	1.69***	1.65***	1.65***	0.68	1.49**	1.14**	0.77	1.29**	1.01***	1.23**
H3<0	-3.83***	-3.66***	-3.94***	-3.63***	-4.63***	-4.07***	-3.84***	-4.03***	-2.82***	-4.06***
H4>0	-1.05	-0.78	-1.16	-2.51	-2.08	-1.30	-1.40	-1.23	-1.42	-1.89
H5>0	2.78***	2.88***	2.78***	1.13*	2.55**	2.78***	2.43***	2.80***	1.39***	2.17**
Observations	819	819	819	815	583	819	819	819	819	607
R^2	0.06	0.05	0.11	0.07	0.09	0.07	0.06	0.07	0.06	0.08

Notes: All columns show the results from estimating equation (4). Sample restricted to respondents in the control group. Column (1) shows the baseline results, for reference. Column (2) drops the region fixed effects. Column (3) controls for the socio-demographic variables. Columns (4)-(5) use alternative measures of right individual. Column (4) defines an individual as "right" based on their ideology (ignoring past vote). Column (5) defines it based on past vote (ignoring ideology). Columns (6)-(9) use alternative measures of good assessment. Column (6) defines good assessment as being below the median prior, regardless of official waiting time. Column (7) defines an individual as good assessment if she is below the median in the difference between the prior and ideal waiting time. Column (8) defines good assessment as being below the median prior, where the median prior is defined separately for each combination of region and individual's ideology. Column (9) considers a continuous variable, defined as the standardized difference in Official - Prior. Column (10) drops Catalonia from the sample. Rows below the constant show results from testing the hypotheses outlined in Section 2, one-sided p-values of the tests are in square brackets. Robust standard errors in parentheses. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$. See notes to Table C.4 for more details.

Table C.14

Divided Regions Controlling for United, Observational Evidence: Robustness

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)
	Base.	No FE	Controls	Right: ideal.	Right: vote	Opt.: median	Opt.: ideal	Opt.: reg.-ideal.	Opt.: cont.	No Catalonia
Good Assessment	0.98 (0.80)	1.12* (0.80)	1.17* (0.79)	1.46** (0.79)	1.10 (0.95)	0.79 (0.81)	0.24 (0.81)	1.27* (0.79)	0.31 (0.42)	0.98 (0.81)
Right-wing Voter	-1.77*** (0.73)	-1.53** (0.73)	-1.65** (0.72)	-1.40** (0.73)	-2.82*** (0.94)	-1.70** (0.85)	-2.08*** (0.86)	-1.84** (0.85)	-2.11*** (0.57)	-1.77*** (0.73)
Good Assessment × Right-wing Voter	-0.60 (1.17)	-0.85 (1.19)	-0.56 (1.16)	-1.49* (1.16)	-0.62 (1.48)	-0.63 (1.15)	0.17 (1.16)	-0.29 (1.14)	-0.51 (0.52)	-0.60 (1.17)
Good Assessment × Divided Region	-2.07** (1.00)	-2.35*** (1.00)	-2.30** (0.99)	-1.91** (1.02)	-2.16** (1.16)	-2.43*** (1.00)	-1.90** (0.99)	-2.78*** (0.97)	-0.70* (0.52)	-1.93** (1.09)
Right-wing Voter × Divided Region	-2.06** (0.93)	-2.13** (0.93)	-2.31*** (0.92)	-2.23*** (0.92)	-1.81* (1.21)	-2.37** (1.06)	-1.76* (1.07)	-2.19** (1.05)	-0.71 (0.73)	-2.29*** (0.98)
Good Assessment × Right-wing Voter × Divided Region	3.37** (1.47)	3.72*** (1.48)	3.36** (1.46)	2.61** (1.45)	3.18** (1.86)	3.41*** (1.45)	2.27* (1.46)	3.09** (1.44)	1.91*** (0.73)	2.77** (1.56)
Constant	2.58*** (0.30)	2.40*** (0.53)	2.57*** (0.30)	2.76*** (0.32)	2.92*** (0.35)	2.79*** (0.34)	2.85*** (0.33)	2.68*** (0.34)	2.43*** (0.23)	2.81*** (0.34)
H1<0	-2.07**	-2.35***	-2.30**	-1.91**	-2.16**	-2.43***	-1.90**	-2.78***	-0.70*	-1.93**
H2>0	1.30	1.38	1.07	0.71	1.01	0.97	0.36	0.31	1.21***	0.84
H3<0	-2.06**	-2.13**	-2.31***	-2.23***	-1.81*	-2.37**	-1.76*	-2.19**	-0.71	-2.29***
H4>0	1.31	1.60*	1.05	0.38	1.37	1.03	0.51	0.90	1.20	0.48
H5>0	3.37**	3.72***	3.36**	2.61**	3.18**	3.41***	2.27*	3.09**	1.91***	2.77**
Observations	1253	1253	1253	1247	873	1253	1253	1253	1253	1041
R ²	0.07	0.05	0.10	0.07	0.09	0.07	0.07	0.07	0.06	0.07

Notes: All columns show the results from estimating equation (5). Sample restricted to respondents in the control group. Column (1) shows the baseline results, for reference. Column (2) drops the region fixed effects. Column (3) controls for the socio-demographic variables. Columns (4)-(5) use alternative measures of right individual. Column (4) defines an individual as "right" based on their ideology (ignoring past vote). Column (5) defines it based on past vote (ignoring ideology). Columns (6)-(9) use alternative measures of good assessment. Column (6) defines good assessment as being below the median prior, regardless of official waiting time. Column (7) defines an individual as good assessment if she is below the median in the difference between the prior and ideal waiting time. Column (8) defines good assessment as being below the median prior, where the median prior is defined separately for each combination of region and individual's ideology. Column (9) considers a continuous variable, defined as the standardized difference in Official - Prior. Column (10) drops Catalonia from the sample. Rows below the constant show results from testing the hypotheses outlined in Section 2, one-sided p-values of the tests are in square brackets. Robust standard errors in parentheses. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$. See notes to Table C.4 for more details.

Table C.15

Divided Regions, Experimental Results: Robustness

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)
	Base.	No FE	Controls	Right: ideol.	Right: vote	Opt.: median	Opt.: ideal	Opt.: reg.-ideol.	Opt.: cont.	No Catalonia
Good Assessment	-1.09** (0.58)	-1.23** (0.59)	-1.19** (0.57)	-0.44 (0.64)	-1.08* (0.66)	-1.66*** (0.57)	-1.65*** (0.57)	-1.49*** (0.57)	-0.38 (0.30)	-0.95* (0.72)
Right-wing Voter	-3.80*** (0.57)	-3.66*** (0.57)	-3.86*** (0.56)	-3.55*** (0.56)	-4.62*** (0.76)	-4.08*** (0.62)	-3.80*** (0.64)	-3.97*** (0.61)	-2.78*** (0.44)	-4.01*** (0.65)
Good Assessment $\times$ Right-wing Voter	2.80*** (0.89)	2.88*** (0.89)	2.83*** (0.88)	1.09 (0.87)	2.59** (1.12)	2.86*** (0.87)	2.42*** (0.88)	2.77*** (0.87)	1.40*** (0.50)	2.21** (1.02)
Treated	0.45 (0.43)	0.49 (0.44)	0.40 (0.43)	0.32 (0.48)	0.39 (0.50)	0.18 (0.47)	0.11 (0.46)	0.19 (0.48)	0.46* (0.34)	0.32 (0.52)
Good Assessment $\times$ Treated	0.01 (0.70)	-0.05 (0.70)	0.17 (0.69)	-0.70 (0.77)	-0.02 (0.79)	0.61 (0.68)	0.84 (0.68)	0.55 (0.68)	-0.05 (0.37)	-0.17 (0.86)
Right-wing Voter $\times$ Treated	0.31 (0.70)	0.25 (0.70)	0.40 (0.69)	0.73 (0.68)	0.92 (0.91)	0.54 (0.76)	0.36 (0.78)	0.54 (0.74)	-0.41 (0.54)	0.43 (0.79)
Good Assessment $\times$ Right-wing Voter $\times$ Treated	-1.88** (1.08)	-1.84** (1.08)	-1.96** (1.07)	-0.57 (1.06)	-2.21* (1.37)	-2.03** (1.07)	-1.75* (1.07)	-2.17** (1.07)	-0.70 (0.59)	-1.05 (1.25)
Constant	2.58*** (0.36)	2.56*** (0.37)	2.61*** (0.36)	2.98*** (0.40)	2.79*** (0.42)	2.94*** (0.40)	2.87*** (0.38)	2.88*** (0.40)	2.15*** (0.28)	2.98*** (0.44)
HT1<0	0.45	0.49	0.40	0.32	0.39	0.18	0.11	0.19	0.46	0.32
HT2>0	0.45	0.45	0.56	-0.38	0.37	0.79*	0.95**	0.75*	0.41	0.15
HT3>0	0.76*	0.74*	0.80*	1.05**	1.31**	0.72	0.47	0.73*	0.05	0.74
HT4<0	-1.12**	-1.15**	-1.00*	-0.22	-0.92	-0.69	-0.44	-0.89*	-0.70	-0.48
HE1>0	0.01	-0.05	0.17	-0.70	-0.02	0.61	0.84	0.55	-0.05	-0.17
HE2<0	-1.87**	-1.89**	-1.80**	-1.27**	-2.23**	-1.41**	-0.90	-1.62**	-0.75*	-1.22*
HE3>0	0.31	0.25	0.40	0.73	0.92	0.54	0.36	0.54	-0.41	0.43
HE4<0	-1.57**	-1.59**	-1.56**	0.16	-1.29	-1.49**	-1.39**	-1.64**	-1.11*	-0.62
HE5<0	-1.88**	-1.84**	-1.96**	-0.57	-2.21*	-2.03**	-1.75*	-2.17**	-0.70	-1.05
Observations	2485	2485	2485	2476	1764	2485	2485	2485	2485	1839
R ²	0.07	0.06	0.09	0.06	0.08	0.07	0.07	0.07	0.07	0.07

Notes: All columns show the results from estimating equation (6). Sample restricted to respondents in divided regions. Column (1) shows the baseline results, for reference. Column (2) drops the region fixed effects. Column (3) controls for the socio-demographic variables. Columns (4)-(5) use alternative measures of right individual. Column (4) defines an individual as "right" based on their ideology (ignoring past vote). Column (5) defines it based on past vote (ignoring ideology). Columns (6)-(9) use alternative measures of good assessment. Column (6) defines good assessment as being below the median prior, regardless of official waiting time. Column (7) defines an individual as good assessment if she is below the median in the difference between the prior and ideal waiting time. Column (8) defines good assessment as being below the median prior, where the median prior is defined separately for each combination of region and individual's ideology. Column (9) considers a continuous variable, defined as the standardized difference in Official - Prior. Column (10) drops Catalonia from the sample. Rows below the constant show results from testing the hypotheses outlined in Section 2, one-sided p-values of the tests are in square brackets. Robust standard errors in parentheses. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$. See notes to Table C.7 for more details.

Table C.16

Divided Regions Controlling for United, Experimental Results: Robustness

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)
	Base.	No FE	Controls	Right: ideol.	Right: vote	Opt.: median	Opt.: ideal	Opt.: reg-ideol.	Opt.: cont.	No Catalonia
Good Assessment	1.01* (0.78)	1.12* (0.80)	1.07* (0.78)	1.42** (0.77)	1.09 (0.90)	0.88 (0.79)	0.36 (0.80)	1.41** (0.78)	0.38 (0.41)	1.01* (0.78)
Right-wing Voter	-1.65** (0.73)	-1.53** (0.73)	-1.57** (0.72)	-1.37** (0.73)	-2.57*** (0.93)	-1.55** (0.85)	-1.88** (0.86)	-1.60** (0.85)	-1.99*** (0.57)	-1.65** (0.73)
Good Assessment $\times$ Right-wing Voter	-0.67 (1.15)	-0.85 (1.18)	-0.46 (1.14)	-1.51* (1.15)	-0.61 (1.42)	-0.74 (1.13)	-0.03 (1.14)	-0.55 (1.13)	-0.54 (0.51)	-0.67 (1.15)
Good Assessment $\times$ Divided Region	-2.10** (0.98)	-2.35*** (0.99)	-2.22** (0.97)	-1.86** (1.00)	-2.17** (1.12)	-2.54*** (0.97)	-2.01** (0.98)	-2.89*** (0.97)	-0.76* (0.51)	-1.96** (1.07)
Right-wing Voter $\times$ Divided Region	-2.15** (0.93)	-2.13** (0.93)	-2.30*** (0.91)	-2.18*** (0.92)	-2.05** (1.20)	-2.53*** (1.05)	-1.92** (1.07)	-2.37** (1.05)	-0.79 (0.72)	-2.37*** (0.98)
Good Assessment $\times$ Right-wing Voter $\times$ Divided Region	3.47*** (1.45)	3.72*** (1.48)	3.23** (1.44)	2.60** (1.44)	3.20** (1.81)	3.59*** (1.43)	2.45** (1.44)	3.32** (1.43)	1.94*** (0.72)	2.88** (1.54)
Treated	0.43 (0.63)	0.63 (0.63)	0.50 (0.61)	1.00* (0.62)	-0.00 (0.74)	0.43 (0.73)	0.22 (0.74)	0.82 (0.74)	-0.43 (0.48)	0.43 (0.63)
Good Assessment $\times$ Treated	-2.35*** (0.98)	-2.49*** (0.99)	-2.48*** (0.97)	-2.86*** (0.95)	-2.50** (1.12)	-1.74** (0.97)	-1.25* (0.97)	-2.47*** (0.97)	-0.22 (0.51)	-2.35*** (0.98)
Right-wing Voter $\times$ Treated	-0.24 (0.88)	-0.59 (0.88)	-0.38 (0.87)	-1.13* (0.88)	0.43 (1.10)	-0.20 (1.02)	-0.14 (1.04)	-0.36 (1.03)	0.36 (0.69)	-0.24 (0.88)
Good Assessment $\times$ Right-wing Voter $\times$ Treated	1.37 (1.41)	1.54 (1.44)	1.33 (1.39)	2.33** (1.40)	1.80 (1.73)	0.98 (1.37)	0.81 (1.38)	1.24 (1.38)	0.62 (0.67)	1.37 (1.41)
Divided Region $\times$ Treated	0.01 (0.76)	-0.13 (0.77)	-0.10 (0.75)	-0.69 (0.79)	0.39 (0.89)	-0.26 (0.87)	-0.11 (0.87)	-0.63 (0.88)	0.89* (0.59)	-0.12 (0.81)
Good Assessment $\times$ Divided Region $\times$ Treated	2.35** (1.20)	2.44** (1.22)	2.61** (1.19)	2.16** (1.23)	2.48** (1.37)	2.35** (1.18)	2.09** (1.19)	3.02*** (1.18)	0.16 (0.63)	2.18** (1.30)
Right-wing Voter $\times$ Divided Region $\times$ Treated	0.55 (1.12)	0.83 (1.12)	0.77 (1.11)	1.86** (1.11)	0.48 (1.43)	0.74 (1.27)	0.49 (1.30)	0.89 (1.27)	-0.77 (0.88)	0.67 (1.18)
Good Assessment $\times$ Right-wing Voter $\times$ Divided Region $\times$ Treated	-3.25** (1.78)	-3.38** (1.80)	-3.24** (1.75)	-2.90** (1.76)	-4.01** (2.21)	-3.01** (1.74)	-2.55* (1.75)	-3.41** (1.74)	-1.31* (0.89)	-2.42* (1.89)
Constant	2.54*** (0.30)	2.40*** (0.53)	2.54*** (0.29)	2.74*** (0.32)	2.88*** (0.35)	2.75*** (0.34)	2.79*** (0.33)	2.61*** (0.34)	2.40*** (0.23)	2.76*** (0.34)
HT1<0	0.01	-0.13	-0.10	-0.69	0.39	-0.26	-0.11	-0.63	0.89	-0.12
HT2>0	2.37***	2.31***	2.50***	1.48*	2.88***	2.10***	1.99***	2.39***	1.05	2.06**
HT3>0	0.56	0.70	0.66	1.18*	0.88	0.48	0.38	0.27	0.12	0.55
HT4<0	-0.33	-0.24	0.03	0.43	-0.65	-0.17	-0.07	-0.12	-1.03	0.31
HE1>0	2.35**	2.44**	2.61**	2.16**	2.48**	2.35**	2.09**	3.02***	0.16	2.18**
HE2<0	-0.90	-0.94	-0.63	-0.74	-1.53	-0.66	-0.46	-0.39	-1.15**	-0.24
HE3>0	0.55	0.83	0.77	1.86**	0.48	0.74	0.49	0.89	-0.77	0.67
HE4<0	-2.70**	-2.55**	-2.47**	-1.04	-3.53**	-2.27**	-2.06**	-2.52**	-2.08**	-1.75
HE5<0	-3.25**	-3.38**	-3.24**	-2.90**	-4.01**	-3.01**	-2.55*	-3.41**	-1.31*	-2.42*
Observations	3806	3806	3806	3793	2673	3806	3806	3806	3806	3160
R ²	0.06	0.05	0.09	0.06	0.08	0.06	0.06	0.06	0.06	0.07

Notes: All columns show the results from estimating equation (8). Column (1) shows the baseline results, for reference. Column (2) drops the region fixed effects. Column (3) controls for the socio-demographic variables. Columns (4)-(5) use alternative measures of right individual. Column (4) defines an individual as "right" based on their ideology (ignoring past vote). Column (5) defines it based on past vote (ignoring ideology). Columns (6)-(9) use alternative measures of good assessment. Column (6) defines good assessment as being below the median prior, regardless of official waiting time. Column (7) defines an individual as good assessment if she is below the median in the difference between the prior and ideal waiting time. Column (8) defines good assessment as being below the median prior, where the median prior is defined separately for each combination of region and individual's ideology. Column (9) considers a continuous variable, defined as the standardized difference in Official - Prior. Column (10) drops Catalonia from the sample. Rows below the constant show results from testing the hypotheses outlined in Section 2, one-sided p-values of the tests are in square brackets. Robust standard errors in parentheses. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$. See notes to Table C.7 for more details.

Table C.17

Attribution of Responsibility: Experimental Evidence, by Treatment Type

	(1)	(2)	(3)	(4)	(5)	(6)
	Div., no rank	United, no rank	Full, no rank	Div., rank	United, rank	Full, rank
Good Assessment	-1.07** (0.59)	1.12* (0.79)	1.12* (0.79)	-1.11** (0.59)	0.87 (0.80)	0.87 (0.79)
Right-wing Voter	-3.82*** (0.57)	-1.67** (0.73)	-1.67** (0.73)	-3.78*** (0.57)	-1.69** (0.74)	-1.69** (0.73)
Good Assessment × Right-wing Voter	2.79*** (0.89)	-0.68 (1.16)	-0.68 (1.16)	2.80*** (0.89)	-0.62 (1.16)	-0.62 (1.16)
Treated	0.60 (0.50)	0.12 (0.72)	0.12 (0.72)	0.30 (0.49)	0.71 (0.72)	0.71 (0.71)
Good Assessment × Treated	-0.10 (0.80)	-3.32*** (1.13)	-3.32*** (1.13)	0.09 (0.79)	-1.16 (1.14)	-1.16 (1.13)
Right-wing Voter × Treated	0.11 (0.84)	0.26 (1.01)	0.26 (1.00)	0.50 (0.79)	-0.64 (1.01)	-0.64 (1.01)
Good Assessment × Right-wing Voter × Treated	-1.82* (1.26)	2.82** (1.63)	2.82** (1.62)	-1.90* (1.24)	-0.50 (1.63)	-0.50 (1.63)
Good Assessment × Divided Region			-2.19** (0.98)			-1.98** (0.99)
Right-wing Voter × Divided Region			-2.15** (0.93)			-2.09** (0.93)
Good Assessment × Right-wing Voter × Divided Region			3.48*** (1.46)			3.42*** (1.46)
Treated × Divided Region			0.49 (0.88)			-0.41 (0.86)
Good Assessment × Treated × Divided Region			3.22*** (1.38)			1.24 (1.39)
Right-wing Voter × Treated × Divided Region			-0.14 (1.31)			1.13 (1.28)
Good Assessment × Right-wing Voter × Treated × Divided Region			-4.64** (2.05)			-1.40 (2.04)
Constant	2.59*** (0.37)	2.45*** (0.53)	2.54*** (0.30)	2.58*** (0.37)	2.51*** (0.53)	2.55*** (0.30)
HT1<0	0.60 [0.89]	0.12 [0.56]	0.49 [0.71]	0.30 [0.73]	0.71 [0.84]	-0.41 [0.32]
HT2>0	0.51 [0.21]	-3.20 [1.00]	3.71 [0.00]	0.38 [0.27]	-0.45 [0.70]	0.83 [0.22]
HT3>0	0.72 [0.14]	0.37 [0.30]	0.34 [0.36]	0.79 [0.10]	0.07 [0.46]	0.72 [0.22]
HT4<0	-1.20 [0.05]	-0.13 [0.45]	-1.07 [0.18]	-1.02 [0.08]	-1.59 [0.04]	0.56 [0.69]
HE1>0	-0.10 [0.55]	-3.32 [1.00]	3.22 [0.01]	0.09 [0.46]	-1.16 [0.84]	1.24 [0.19]
HE2<0	-1.92 [0.02]	-0.50 [0.33]	-1.42 [0.17]	-1.82 [0.03]	-1.66 [0.08]	-0.16 [0.46]
HE3>0	0.11 [0.45]	0.26 [0.40]	-0.14 [0.54]	0.50 [0.26]	-0.64 [0.73]	1.13 [0.19]
HE4<0	-1.71 [0.04]	3.07 [0.99]	-4.78 [0.00]	-1.40 [0.07]	-1.14 [0.19]	-0.27 [0.43]
HE5<0	-1.82 [0.07]	2.82 [0.96]	-4.64 [0.01]	-1.90 [0.06]	-0.50 [0.38]	-1.40 [0.25]
Observations	1636	895	2531	1668	860	2528
R ²	0.07	0.05	0.07	0.06	0.08	0.07

Notes: Columns (1)-(3) restrict the sample to respondents either in the control group or in the treatment group not receiving the additional information about how the relevant waiting time in their region ranks among all regions in Spain. Columns (4)-(6) restrict the sample to respondents either in the control group or in the treatment group receiving the additional information about how the relevant waiting time in their region ranks among all regions in Spain. Columns (1), (2), (4), and (5) show the results from estimating equation (6) for respondents in regions with divided and united governments, respectively. Columns (3) and (6) show the results from estimating equation (8) for the full sample. A region is classified as "divided" if none of the parties in the central government coalition participates in the regional government coalition; see Table C.1 for a list of divided regions. The dependent variable is how much responsibility the respondent assigns to the regional versus the central government, measured on a scale from -10 (all responsibility lies with the central government) to 10 (all responsibility lies with the regional government). "Left-wing" and "Right-wing" refer to respondents' partisan affiliation based on past regional vote or ideological self-placement. "Bad Assessment" and "Good Assessment" indicate whether respondents' prior beliefs about waiting times exceed or fall below the official regional average, respectively. Rows below the constant show results from testing the hypotheses outlined in Section 2, one-sided p-values of the tests are in square brackets. Robust standard errors in parentheses. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

D Differences from the Pre-Analysis Plan

The survey and experiment analyzed in this paper were pre-specified in a pre-analysis plan (PaP) that we registered with the AEA Randomized Controlled Trials Registry on 7 July 2023. We also obtained approval for the survey and experimental design from the ethics committee at CEMFI (Application Reference Number 3; Approval date: 18 June 2023). This Appendix describes how the analyses reported in the paper relate to that plan.

The pre-analysis plan was written to cover a broad research agenda spanning several projects that draw on the same survey. It therefore pre-specifies a large set of outcomes, hypotheses, and specifications, organized into six families of outcomes—beliefs about the provision and governance of public healthcare; trust in political institutions; political preferences and support for the incumbent governments; polarization; support for taxation and redistribution; and perceived health status and well-being—together with a range of heterogeneity analyses and robustness checks.

This paper develops and tests one component of that agenda: a model of attribution of responsibility under ambiguity and its empirical implications. Accordingly, we report only the subset of pre-registered analyses that speaks to this theory—principally beliefs about which layer of government is responsible for public healthcare, and how those beliefs respond to information about performance—and we do not report results for the remaining pre-registered outcome families. Those additional analyses are outside the scope of the present paper and have not been carried out here; we would be glad to conduct and share them upon request.

Within the subset of analyses we do report, our empirical implementation follows the pre-analysis plan. As described in Section 3, we apply the pre-specified sample-selection rules, dropping respondents who fail the attention check or complete the questionnaire in an implausibly short time, as well as a single observation with an extreme outlier in reported waiting-time beliefs. We classify partisanship, assessment of the quality of public healthcare, and attribution of responsibility using the definitions set out in the plan, and we estimate the pre-registered experimental and observational specifications. Where the exposition departs in minor ways from the exact wording of the plan—for example, in the labeling of variables or the grouping of specifications into tables—these differences are cosmetic and do not affect the results.

E Questionnaire

This appendix reproduces the respondent-facing survey items used to construct variables, implement the information treatments, define the analysis sample, or describe the survey in the paper. Questions not used for those purposes and programming instructions are omitted. Bracketed placeholders were filled with respondent-specific information. Panel-profile variables used in the paper but not asked in this questionnaire (such as sex and municipality type) are not reproduced. The order of some modules and response options was randomized. Questions translated from Spanish by the authors.

E.1 Background characteristics and political preferences

- **Region of residence.** In which autonomous community do you live?

Response: Andalusia; Aragon; Cantabria; Castile and Leon; Castile–La Mancha; Catalonia; Ceuta; Community of Madrid; Community of Navarre; Valencian Community; Extremadura; Galicia; Balearic Islands; Canary Islands; La Rioja; Melilla; Basque Country; Principality of Asturias; or Region of Murcia.

- **Education.** What is the highest level of education you have completed?

Response: No formal education; primary education; lower-secondary education; upper-secondary education; first- or second-level vocational education; unfinished university studies; three-, four-, or five-year university degree; master’s degree; PhD; other; or prefer not to answer.

- **Age.** How old are you?

Response: Open numerical response.

- **Employment status.** Which of the following best describes your current employment situation?

Response: Employee in the private sector; employee in the public sector; business owner, professional, or self-employed; retired or pensioner; unemployed but previously employed; unemployed and never employed; student; unpaid domestic work; or other.

- **Household composition.** How many people, including yourself, live in your household? Please enter the number of people in each age group.

Response: Children aged 0–17; adults aged 18–64; and adults aged 65 or older.

- **Participation in the previous regional election.** Did you vote in the most recent regional election held in your autonomous community? If your community did not hold a regional election on May 28, please answer with reference to the most recent regional election.

Response: Yes; I think I voted; I think I did not vote; no; or do not know/prefer not to answer.

- **Past regional vote.** Which party did you vote for in the previous regional election? If your community did not hold a regional election on May 28, please answer with reference to the most recent regional election.

Response: The list of parties standing in the respondent’s autonomous community, plus other, abstention/did not vote, and do not know.

- **Left–right self-placement.** In politics, people commonly use the terms left and right. On a scale from 1 to 10, where 1 means “Left” and 10 means “Right,” where would you place yourself?

Response: 1 (Left) to 10 (Right).

- **Attention check.** Recent studies of decision-making show that decisions are affected by the context in which they are made and reflect people’s prior feelings, knowledge, experience, and environment. Thank you for helping make the survey results meaningful by following the instructions. To help us confirm that you have read these instructions, select “none of the above” from the following options.

Response: Anger; joy; sadness; fear; surprise; or none of the above.

E.2 Waiting-time beliefs and experimental information

Respondents were randomly assigned to begin with either the specialist-consultation module or the elective-surgery module. Treated respondents received official information about the waiting-time dimension asked first. The other waiting-time dimension was elicited later. Control respondents did not receive the official information. The two treatment variants differed in whether respondents also saw their autonomous community's position in the national ranking.

E.2.1 Specialist consultations

- **Ideal waiting time.** In your opinion, what would be an appropriate waiting time to see a specialist doctor in the public healthcare system, averaged across medical specialties?

Move the slider to select your preferred waiting time. Please bear in mind that reducing waiting times could require higher taxes or fewer resources for other public services.

Response: 0 to 180 days, where 180 denotes 180 days (six months) or more.

Note: The specialties considered were gynecology, ophthalmology, traumatology, dermatology, otorhinolaryngology, neurology, general surgery, urology, gastroenterology, and cardiology.

- **Belief about the actual waiting time.** We would now like you to try to guess how many days patients have to wait to see a specialist. Can you tell us how many days patients in [Autonomous Community] had to wait in 2022 to see a specialist doctor in the public healthcare system, averaged across medical specialties? Move the slider to indicate your answer.

Response: 0 to 180 days, where 180 denotes 180 days (six months) or more. Respondents selecting 180 or more were asked to enter the exact number of days.

- **Belief about relative waiting times.** How do you think waiting times to see a specialist doctor in the public healthcare system in [Autonomous Community] compare with the average across all autonomous communities?

Response: Much shorter than average; shorter than average; approximately average; longer than average; or much longer than average.

E.2.2 Elective surgery

- **Ideal waiting time.** In your opinion, what would be an appropriate waiting time for surgery in the public healthcare system, averaged across different types of surgery?

Move the slider to select your preferred waiting time. Please bear in mind that reducing waiting times could require higher taxes or fewer resources for other public services.

Response: 0 to 180 days, where 180 denotes 180 days (six months) or more.

Note: The procedures considered were cataract surgery, inguinal/femoral hernia repair, hip replacement, arthroscopy, lower-limb varicose-vein surgery, cholecystectomy, hallux valgus surgery, adenotonsillectomy, benign prostatic hyperplasia surgery, pilonidal-cyst surgery, carpal-tunnel surgery, knee replacement, valvular heart surgery, coronary bypass surgery, and hysterectomy.

- **Belief about the actual waiting time.** We would now like you to try to guess how many days patients have to wait for surgery. Can you tell us how many days patients in [**Autonomous Community**] had to wait in 2022 for surgery in the public healthcare system, averaged across different types of surgery? Move the slider to indicate your answer.

Response: 0 to 180 days, where 180 denotes 180 days (six months) or more. Respondents selecting 180 or more were asked to enter the exact number of days.

- **Belief about relative waiting times.** How do you think waiting times for surgery in the public healthcare system in [**Autonomous Community**] compare with the average across all autonomous communities?

Response: Much shorter than average; shorter than average; approximately average; longer than average; or much longer than average.

E.2.3 Information shown to treated respondents

Official waiting-time information. Respondents assigned to treatment saw the corresponding message for the dimension elicited first:

According to official public data for 2022, in [**Autonomous Community**] the average waiting time to see a specialist doctor in the public healthcare system is [XX] days. This is [ZZZ] days [**more/less**] than you predicted.

or

According to official public data for 2022, in [**Autonomous Community**] the waiting time for surgery in the public healthcare system is [XX] days. This is [ZZZ] days [**more/less**] than you predicted.

Additional ranking information. Half of the treated respondents also saw the corresponding regional-ranking message:

We now show you the ranking of autonomous communities from the shortest to the longest waiting time to see a specialist doctor in the public healthcare system. According to official public data, [**Autonomous Community**] has [**ranking description**] waiting time, and its waiting time is [**above/below**] the national average. You had predicted that your community was [**respondent's relative-waiting-time answer**].

or

We now show you the ranking of autonomous communities from the shortest to the longest waiting time for surgery in the public healthcare system. According to official public data, [**Autonomous Community**] has [**ranking description**] waiting time, and its waiting time is [**above/below**] the national average. You had predicted that your community was [**respondent's relative-waiting-time answer**].

E.2.4 Belief elicited after the information module

- Respondents who first answered the specialist-consultation questions were later asked: We would now like you to try to guess how many days patients have to wait for surgery. Can you tell us how many days patients in [**Autonomous Community**] had to wait in 2022 for surgery in the public healthcare system, averaged across different types of surgery? Move the slider to indicate your answer.

Response: 0 to 180 days, where 180 denotes 180 days (six months) or more. Respondents selecting 180 or more were asked to enter the exact number of days.

- Respondents who first answered the elective-surgery questions were later asked: We would now like you to try to guess how many days patients have to wait to see a specialist. Can you tell us how many days patients in [**Autonomous Community**] had to wait in 2022 to see a specialist doctor in the public healthcare system, averaged across medical specialties? Move the slider to indicate your answer.

Response: 0 to 180 days, where 180 denotes 180 days (six months) or more. Respondents selecting 180 or more were asked to enter the exact number of days.

E.3 Attribution of responsibility

- **Political (de facto) responsibility.** We would like to ask which layer of government you believe is more responsible for the quality of public healthcare. On a scale where -10 means “the Government of Spain bears all the responsibility” and 10 means “the government of your autonomous community bears all the responsibility,” how much responsibility would you attribute to each government?

Response: Slider from -10 (Government of Spain) to 10 (government of the respondent’s autonomous community).

- **Legal (de jure) responsibility.** Which public administration has authority over healthcare—that is, who decides the level of spending, personnel management, hospital management,

primary care, and so forth?

Response: The central government; the government of your autonomous community; your municipal government; the institutions of the European Union; or do not know/prefer not to answer.