

Positional Convergence and Legislative Success: Business Lobbying in EU Public Consultations

Gabriel Rodríguez Molina

Universitat Autònoma de Barcelona — Department of Political Science and Public Law

Draft paper. Please do not cite or circulate without the author's permission.

Abstract

Who wins and who loses among business stakeholders in EU legislative politics? While business interests dominate EU public consultations, what determines whose policy preferences are attained in the final legislative text remains an open question. This paper argues that lobbying success at policy legitimization, the stage at which stakeholders comment on a proposal the Commission has formulated and adopted but the co-legislators have not yet amended, is driven by signalling, so that the relevant property is how widely a stakeholder's preferences are shared across the field, which indicates to policymakers where support for a proposal lies. Using an original dataset of 1,908 business stakeholder–policy observations across 44 EU legislative files from the 2019–2024 legislature, I measure preference attainment at the level of the individual issue a stakeholder engaged, and construct a signed network measure of positional convergence using sentence embeddings and large language model classification. Stakeholders whose demands are widely shared are more successful, whereas the resources an organisation commands and the access it enjoys bear no relation to attainment, nor do the complexity of the file and its salience. The effect is amplified in polarised consultations, where the field divides into mutually opposed groupings, and stakeholders whose positions diverge from the prevailing regulatory orientation attain markedly less. These findings suggest that preference attainment in EU consultations is shaped by where an actor's preferences sit in the aggregate distribution of preferences rather than by organisational endowment.

Keywords: lobbying success, positional convergence, signed networks, EU public consultations, preference attainment, policy space, NLP

1. Introduction

The European Union's legislative process unfolds under a persistent democratic tension. While EU institutions have progressively expanded channels for stakeholder input, for example through the Have Your Say public consultation portal, business stakeholders (trade associations, companies, and corporate federations) dominate these consultations, and what determines who among them succeeds remains an open question (Saurugger 2008; Schoenefeld 2021). The submissions analysed in this study are drawn from policy legitimisation, the feedback period that opens once the Commission has formulated and adopted a legislative proposal and closes before the co-legislators begin amending it. Stakeholders at this point address a finished legal instrument rather than an open agenda, and their submissions state positions on provisions that already exist and remain open to amendment. Research on who succeeds in EU consultations has focused largely on policy formulation, where stakeholders respond to a proposal still under preparation (Bunea 2013; Chalmers 2020). Success at legitimisation remains largely unexamined, despite evidence that participation at this stage covaries with the extent to which the proposal is subsequently modified (Bunea & Lipcean 2026).

The lobbying literature has long sought to explain lobbying outcomes through individual actor characteristics such as resources, expertise and institutional access (Bouwen 2002, 2004; Eising 2007). Yet the empirical record on whether these variables predict whose preferences reach the final legislative text remains inconclusive (Mahoney 2007a; Dür et al. 2015), suggesting that lobbying is better understood as a collective process shaped by the policy context in which demands are advanced (Mahoney 2007a). Work in this vein has established that collective properties predict outcomes where individual traits do not, whether aggregate information supply and market power (Klüver 2012, 2013), the diversity of actor types within a coalition (Junk 2019), or industry cohesion (Chalmers 2020). It has done so, however, by taking the grouping as the unit of analysis, whether identified in advance from formal association membership or reported cooperation (Beyers & De Bruycker 2018) or recovered from the positions actors state and reduced to the camp each actor belongs to (Baumgartner et al. 2009; Klüver 2012, 2013), so that every member is assigned the same value. The approach taken here build on the collective approach on lobbying while locating each actor individually within the field of stated positions, recording how far its demands coincide with those of others addressing the same provisions.

This paper asks where each actor's demands sit within the field of positions a consultation produces. Drawing on network theory (Freeman 1978/79; Borgatti 2005) and on studies that treat the collective expression of preferences as a signal to decision-makers of where support for a proposal lies (Mahoney 2007b; Nelson & Yackee 2012; Phinney 2017; Junk 2019), I develop an account in which lobbying success is measured by tracking how the text a stakeholder engaged changes between the Commission proposal and the adopted regulation. That outcome is explained by the stakeholder's own location in the field of positions, captured by *positional convergence*, a signed network measure of how far its demands coincide with, or run against, those of the other actors engaging the same provisions.

The analysis covers the European Commission's Have Your Say consultation process across 44 EU legislative files from the 2019–2024 legislature, comprising 1,908 business stakeholder–policy observations. Both measures are constructed through a natural language processing (NLP) method combining text extraction, semantic matching, and large language model classification. The findings show that neither organisation-level characteristics (financial resources, staff size, institutional access, umbrella membership, and the length of the submission filed) nor policy-level ones (the legislative complexity of a file and the breadth of participation it attracts) predict whose preferences prevail. Positional convergence does, and the effect is amplified in polarised consultations, where the field divides into mutually opposed groupings. Besides, business stakeholders whose positions diverge from the prevailing regulatory orientation attain markedly less, and positional convergence appears to matter more for them than for aligned stakeholders, though this last association is not statistically significant.

The contribution of this paper is to explain lobbying success through positional convergence, extending to policy legitimation an insight established at formulation, that an actor's prospects depend on where its preferences sit relative to those of others (Klüver 2011; Bunea 2013). Carrying that insight across stages requires measuring the positional property differently. At formulation, the Commission poses questions, and the issues on which positions are compared can be read off the consultation document (Bunea & Ibenskas 2017; Chalmers 2020) or the items of a questionnaire (Bunea, Wüest & Lipcean 2025). At legitimation there is only the Commission's proposal, a finished legal text, so the issues cannot be inherited from the instrument that elicited the responses. They are instead derived from the text itself and inferred through engagement, since issues are not given but issued, an article becoming one only when stakeholders address it.

A second alternative design choice concerns the substantive content of a demand, since existing designs place demands on a scale running from less to more regulation, which presumes that what actors ask for differs in degree. This study argues that a scale is less suited to demands framed against a drafted text, where a stakeholder proposing an alternative compliance mechanism asks for neither more nor less regulation, and two such proposals may be mutually exclusive without either being the more stringent. The framework instead records what a demand seeks, classifying whether a stakeholder accepts a provision, seeks its removal, or seeks to replace it with a different rule, which preserves the difference between alternatives that are genuinely incompatible.

Section 2 reviews what explains lobbying success and what is known about success within public consultation procedures. Section 3 develops the theoretical framework and derives the hypotheses. Section 4 describes the research design and methodology, including data construction, variable operationalisation, and estimation strategy. Section 5 presents the results. Section 6 discusses implications and limitations. Section 7 concludes.

2. Literature Review

This section reviews two streams of research. The first concerns what explains lobbying success, moving from accounts resting on the attributes of individual actors to those treating lobbying as a collective enterprise, and distinguishing within the latter between coalitions formed by deliberate cooperation and sides identified from the positions actors take, on both of which the present study draws. The second concerns the evidence on success within public consultation procedures, distinguishing studies of policy formulation from those following the text through to adoption, which establishes and justifies the stage this paper analyses.

2.1 Explaining Lobbying Success

Following common practice in this literature, the outcome analysed is framed in terms of success rather than influence, a distinction that turns on causal attribution, since success records whether an actor obtained what it sought while influence claims that the actor produced the outcome (Dür 2008; Klüver 2013; Binderkrantz & Pedersen 2019). Success so understood is measured as preference attainment, the correspondence between the position an actor stated and the content of the policy finally adopted, which the literature treats as the standard operationalisation of the concept. Attributing policy change to particular actors would require observing pathways that the opacity of informal exchange makes inaccessible (Woll 2007; Bell & Hindmoor 2009).

In this regard, a first line of work explains success by the attributes of individual actors, mainly resources, expertise, and institutional access. Its dominant formulation treats attainment as a return on information supplied, since policymakers face technical problems under constraints of time and expertise, interest groups hold information they lack, and access is obtained by supplying it, so that the actors best placed to succeed are those with the staff, budget and technical capacity to supply expertise of the kind policymakers require (Bouwen 2002, 2004; Eising 2007; Dür & Mateo 2016). The empirical record on whether these attributes predict whose preferences reach the final text nonetheless remains inconclusive (Mahoney 2007a; Dür et al. 2015), with policy-level context often dominating actor characteristics (Mahoney 2007a).

A second line of work treats lobbying as a collective enterprise, analysing coalitions formed through deliberate, strategic cooperation, such as umbrella associations or coordinated alliances (Hula 1999; Beyers & De Bruycker 2018). In this regard, De Bruycker & Beyers (2019) find that outside lobbying is more effective

when combined with coalition membership, and Junk (2019) that diversity within a coalition broadens the range of arguments and constituencies a demand carries. Others in this second line, rather than studying formal coalitions, identify groupings *ex post* from the positions actors actually take, treating them as emergent patterns of coincident demands rather than pre-existing structures (Klüver 2012, 2013). In this regard, if coalitions are given in advance and actors are compared with the organisation they belong to (Beyers & De Bruycker 2018; Junk 2019), sides and blocs are identified from stated positions (Baumgartner et al. 2009; Atikcan & Chalmers 2019), so that any two actors addressing the same matter can be compared, whether or not an organisational relation connects them, and the structure of alignment and opposition across a consultation becomes observable. This paper draws on both strands, taking from the coalition literature the mechanism by which collective demands operate, considering that the distribution of positions is a signal sent to policymakers (Mahoney 2007b; Nelson & Yackee 2012; Phinney 2017); while locating that signal in the positions actors express rather than in the organisations connecting them. This paper likewise treats groupings as emergent from stated positions, departing from the sides literature only in declining to fix discrete blocs, and characterises each actor by its position within the resulting structure of preferences.

2.2 Lobbying Success in Public Consultation Procedures

The Better Regulation framework provides three distinct opportunities to participate throughout the policy-making process: a call for evidence preceding drafting, a public consultation during a proposal's development structured by a questionnaire, and a feedback period following the College's adoption of the proposal, carrying no questionnaire (European Commission 2021, Boxes 4 and 6). The first two belong to policy formulation and the third to policy legitimisation (Bunea & Lipcean 2026).

In this regard, studies analysing the proposal ask whose preferences reach the Commission's draft. Bunea (2013) codes positions from free-text submissions in environmental policy and finds that what predicts attainment is where an actor's demand sits in the distribution of demands made by others, a group holding the median position relative to others being substantially more likely to obtain it, an effect she reads as evidence of the consensual character of EU policymaking (Downs 1957; Thomson 2009). Chalmers (2020) reaches a compatible conclusion for financial-sector consultations, where success depends on industry unity rather than

the resources of individual firms. Both establish that the policy environment matters more than the attributes of participants.

Designs of this kind predominate in the literature (Bunea & Ibenskas 2017; Chalmers 2020; Bunea, Wüest & Lipcean 2025), since the Commission states that formulation-stage consultations serve to gather the input on which it draws while drafting, so that positions filed there have a declared route into the text. No comparable statement covers what the Parliament and Council do with feedback filed once a proposal has been adopted, since the co-legislators are under no duty to consider material gathered by another institution (European Commission 2021; European Court of Auditors 2019), and the stage has accordingly attracted little attention. Nevertheless, some studies carry the analysis beyond the Commission's text to the regulation finally adopted. In this regard, Klüver (2012, 2013) asks whether the adopted act moved closer to a given stakeholder's position, drawing those positions from submissions filed before the Commission puts its proposal forward. Outcomes settled by the Parliament and the Council are thus inferred from positions stated in a Commission consultation, which offers a precedent for this study. Bunea & Lipcean (2026) address that question directly and show that the stage at which preferences are stated matters. Modelling the textual similarity between 315 Commission proposals and the acts subsequently adopted, they find that participation during policy legitimization, after the College has adopted a proposal, covaries with the extent to which the proposal is modified before adoption, whereas participation during formulation does not do so in the same way, and stakeholder support expressed at legitimization is associated with a proposal surviving more nearly intact. Stakeholder positions stated at legitimization thus stand in a measurable relation to subsequent amendment, it is at this stage that the present study locates its question of whose preferences those amendments reflect. Nevertheless, while their outcome of interest is the extent to which the proposal changes, the present paper examines what explains whose preferences are obtained.

Previous work on preference attainment either breaks a policy into the separate issues it raises and records an outcome on each (Bunea 2013; Chalmers 2020), or treats the file as a single object and asks how far it moved towards a stakeholder's position overall (Klüver 2012, 2013). In designs of the former kind, the issue structure is supplied by the Commission, whether by the text of the consultation call (Bunea & Ibenskas 2017), the policy questions put to respondents (Chalmers 2020), or the items of a closed questionnaire (Bunea, Wüest & Lipcean 2025). At legitimization the Commission poses no questions, so the issue structure cannot be taken from a

consultation document. The option preferred here is to recover it from the submissions themselves, an approach better suited to feedback filed without prompts. This study therefore infers issues from stakeholder engagement with the proposal, each article constituting a potential issue by virtue of the differentiated subject matter that drafting convention requires it to address.

3. Theoretical Framework

3.1 Success as an Issue-Level Outcome

Success is measured here as relative measure, for each provision a stakeholder engaged, its stated position is set against what happened to that provision between the Commission proposal and the adopted regulation, scored positively where the outcome matched the position and negatively where it ran against it, and aggregated to the issues the stakeholder engaged. Section 4.2 sets out the scoring in full.

Relative success is assessed at the level of the issue rather than of the policy as a whole. This study assumes that stakeholders engage a proposal selectively, and that the demands an actor advances differ across its parts. A trade association may support Articles 3 and 5 of a draft regulation while proposing alternatives to Article 7 and opposing Article 12, and each of those preferences is resolved separately in the adopted text. Whether such a stakeholder "succeeded" is therefore not a single fact about the file but an aggregate of what happened on each issue it engaged, which allows the same stakeholder to win on some issues and lose on others, making partial success an observable outcome rather than an averaging artefact.

Aggregating outcomes this way requires the file to be divided into the discrete matters on which stakeholders take positions. A policy issue is a body of substantive content on which a preference can be held and about which actors can agree or disagree. The issue is the standard unit of analysis in EU lobbying literature, and it is typically drawn from the document by which the institution structured the consultation. In this regard, Chalmers (2020) maps comment letters onto the policy questions the Commission poses, and Bunea, Wüest & Lipcean (2025) take their issues from the items of a closed questionnaire. The legitimisation stage provides no equivalent document, so the issue structure is derived from the proposal itself, which this framework accomplishes by identifying issues with articles.

Articles are an appropriate unit at which to observe policy issues because drafting convention requires that each article deal with one matter, and articles carry titles naming it, so their thematic differentiation is a property drafters were obliged to build in. Articles are also the unit in which the contest over a file is conducted, since amendments are tabled and voted article by article, and the unit stakeholders themselves use when citing the text. In this regard, policy issues are not given but issued, meaning that an article becomes an issue through stakeholder engagement with it. While the Commission's questionnaire prompts every respondent on every item, recording positions on matters an actor may have no stake in, a stakeholder

writing without such prompts addresses only the provisions that bear on its interests. The number of actors addressing a given article accordingly constitutes an observable indicator of its salience, established by the submissions themselves rather than fixed in advance by the consultation's design. A partial precedent for identifying issues this way is Bunea (2013), who draws them from the consultation call, the questionnaire and the issues the groups themselves raise recurrently across contributions as a third source.

Two further features respond to the characteristics of the legitimization stage, and distinguish how success is analysed here from standard practice in studies of preference attainment in public consultations. Stakeholders' preferences are expressed about particular obligations, thresholds, definitions and exceptions. Depending on how specific a stakeholder's comment is, it may point to a section, a subsection or an independent paragraph within an article. The design is sensitive to these differences of specificity, while comments directed at the initiative as a whole, at the Commission's general approach, or at broad regulatory goals are excluded since they cannot be attached to any identifiable policy issue. A stakeholder demand directed at a concrete provision is considered whether or not it explicitly cites an article number, these article subunits being referred to as *provisions*. Because provisions sit inside one another, similar consultation demands can be matched to nested provisions. A stakeholder may reiterate or develop a demand in different passages of its submission, so that one preference is recorded more than once. Nested provisions are therefore reduced to one before the article is scored, so that no preference is counted twice. The procedure is set out in Section 4.2 and Appendix A.3.

The second feature concerns how the relation between a stated preference and a legal provision is classified. This study holds that recording success by the content of the demand, rather than by its location on a scale of regulatory stringency, is better suited to open comment consultations conducted on a drafted text. Success therefore depends on whether a provision changed in the way a stakeholder posed it, not merely on whether it favours more or less regulation. A position is *aligned* when the stakeholder accepts the provision as drafted; *opposed* when it asks for the provision to be removed; and *divergent* when it proposes alternative framing, whether a narrower scope, a longer deadline, or an alternative compliance mechanism. Two divergent demands on the same provision may be mutually incompatible without either being the more or less stringent, and recording substantive content rather than degree is what allows that incompatibility to be registered rather than resolved onto a scale.

Although both features can be read as matters of measurement accuracy at this stage, each rests on a theoretical claim about the preferences being measured. The first is that no level of the text is essentially the proper one at which preferences are expressed. Legal content is hierarchical and nested, and a demand may bear on a question as narrow as a single threshold or as broad as an entire regime, so the level at which a preference is stated is a property of the preference rather than a parameter for the analyst to fix. The second is that preferences are frequently irreducible to one another. Two stakeholders may both accept that a matter should be regulated and still propose mutually exclusive alternatives, neither of which is the more stringent, so that the distance between them is no smaller than that separating either from an actor seeking the provision's removal. Treating disagreement as distance along a regulatory scale therefore registers as near-agreement what are in substance incompatible demands.

3.2 Positional Convergence and the Signalling Mechanism

This study builds on accounts that treat lobbying as *signalling*, on which group activity indicates to policymakers the breadth and depth of political support a proposal commands (Kingdon 1995; Kollman 1998), and on the work applying that logic to coalitions (Mahoney 2007b; Nelson & Yackee 2012; Phinney 2017). Developed on US congressional lobbying and agency rulemaking, this literature has studied coalitions in terms of actors joined by deliberate cooperation, whether an alliance formed around a particular file or an association standing for a common interest. Policymakers, on this account, "search for clear political signals regarding the overall strength of support or opposition to a policy proposal," and broad coalitions allow advocates to signal "policy-makers that a large majority of the electorate will likely support them" (Mahoney 2007b).

Part of what this work identifies is a property not of the formal organisation but of the content a grouping supports. Nelson & Yackee (2012) find that coalition lobbying shifts policy when a coalition is larger and when there is "consensus across the messages" its members send, the second of which is a property of the positions advanced rather than of the arrangement among those advancing them. Phinney (2017) treats coalitions as addressing "legislators' uncertainty regarding the consequences of their policy decisions", and locates their contribution in which actors are found to hold the same position, since agreement among actors that would ordinarily be opposed cannot be attributed to the interest of any one of them.

In this regard, lobbying research marks the distinction between organised cooperation and convergent preferences. *Coalition* denotes deliberate and strategic cooperation, established from what groups report in interviews or from formal association membership (Hula 1999; Beyers & De Bruycker 2018; De Bruycker & Beyers 2019; Junk 2019; Chalmers 2020), whereas *side* denotes a set of actors adopting the same position, independently of any formal tie between them (Baumgartner et al. 2009; Atikcan & Chalmers 2019). This study builds on the second notion while retaining the signalling mechanism, arguing that what reaches policymakers is the distribution of positions rather than the formal arrangement among those holding them. Coincident positions therefore signal support even where no organisational tie links those expressing them.

This study departs from the sides literature in one respect, since sides as ordinarily conceived place each actor in one camp for the file as a whole (Baumgartner et al. 2009; Atikcan & Chalmers 2019), whereas the framework advanced here imposes no such division. The policy space of an EU consultation is not always organised by a single dimension, some being structured by more than one (Bunea & Ibenskas 2015; Bunea, Wüest & Lipcean 2025), and alignment is not constant within a file, since the same two actors may converge on one provision and divide on the next. The analysis accordingly considers each actor within the structure of agreement and conflict a consultation produces. The main variable of this study, positional convergence, records where an actor stands in that structure, measuring how far its demands are shared or contested by others addressing the same provisions. Measuring position at the level of the actor rather than the grouping has a precedent in Bunea (2013), whose analysis of preference attainment asks whether an actor's position occupies a median location relative to the positions others have taken, and finds this a better predictor than the resources or the formal ties of the group itself.

Positional convergence is expected to predict success because the distribution of positions in which an actor sits reaches the co-legislators, and this expectation rests on two mechanisms. The first concerns institutional transmission: since the submissions analysed are stakeholder texts from the legitimisation stage, the feedback period that opens once the Commission has formulated and adopted a proposal and closes before the co-legislators begin amending it. Feedback filed at this stage is addressed not only to the Commission which collects it, but to the European Parliament and the Council, which will amend the text. The Commission states that at this stage feedback is "summarized by the Commission and presented to the EP and Council, with the aim to feed these views into the legislative debates" (European Commission 2017, p. 439);

while submissions are published on arrival and shared with the co-legislators (European Court of Auditors 2019). What reaches the co-legislators is therefore an aggregate of positions rather than a set of individual representations, which is the form the signalling account requires, and an actor's place within that aggregate is what positional convergence measures.

The second mechanism does not depend on transmission, it rests on the assumption that publicly expressed preferences, while not equivalent to an actor's full strategic agenda, constitute a reliable indicator of the specific policy amendments each stakeholder will pursue through all available channels, including informal consultations, bilateral meetings, and other venues that remain unobservable to the researcher. On this perspective, a submission is a trace of a position rather than the vehicle by which it travels, so that the consultation record maps the demand structure policymakers encounter however that structure reaches them. This is the logic Klüver (2009, p. 536) identifies in noting that preference attainment "covers all channels of influence", since a correspondence between what an actor sought and what the adopted text provides registers the outcome whichever channel was operative.

This structure of expressed preferences is mapped as a network for each policy. Where two stakeholders address the same provision, their positions are compared and recorded as agreement or disagreement, and the balance of these across the provisions they both address determines whether the two actors stand together or in conflict. Positional convergence aggregates those relations for each actor, so that agreement with others raises its score and conflict lowers it. This is a network of similarity rather than of flow in Borgatti's (2005) sense, of the same kind as co-voting or co-citation structures. A tie records a correspondence between two independently stated positions, not a transfer between the actors holding them, and centrality over it is read accordingly. A stakeholder whose demands coincide with those of many others, and attract few opposing positions, is one whose position the co-legislators can accommodate at low political cost. The framework's central expectation follows:

H1 (Positional Convergence): Positional convergence is positively associated with relative success.

Because the feedback period opens only once the College has adopted the proposal (European Commission 2021, Box 6), the text stakeholders address is not a neutral object about which positions are taken but a position the Commission has committed itself to, and which it continues to defend through the negotiation that follows. The *prevailing regulatory orientation* of a file is the position embodied in that tabled proposal. The proposal is therefore not one position among those

expressed but the baseline against which the others are read, which places aligned and divergent stakeholders asymmetrically. A stakeholder aligned with a provision asks for what the text already provides, and its demand is satisfied by inaction; a divergent stakeholder asks the co-legislators to displace a position the Commission is committed to defending. Since most provisions survive unchanged into the adopted text, the aligned stakeholder's demand is met by default, while the divergent stakeholder's requires a change in the direction sought, which few provisions undergo. Existing evidence points the same way, Bunea (2013) finding that status-quo demands attain more than demands for change, a pattern the broader literature attributes to the general advantage defenders of an existing position enjoy (Baumgartner et al. 2009).

H2a (Divergence Penalty): Divergence from the prevailing regulatory orientation is negatively associated with relative success, relative to alignment with it.

Following the framework, the two expectations should interact. A divergent demand requires the co-legislators to ammend text as tabled, and whether they do so depends on how widely this demand is shared across the field. An isolated divergent position gives policy-makers no reason to amend, while one shared across many independently filed submissions signals that the alternative commands support and lowers the political cost of departing from the proposal (Nelson & Yackee 2012).

H2b (Conditional Convergence): The association between positional convergence and relative success is stronger the greater the divergence from the prevailing regulatory orientation, relative to alignment with it.

Nelson & Yackee (2012) make the signal's effect conditional on its clarity, since what officials attend to is a clear indication of where support lies. Where a policy space is weakly divided, positions neither consolidate into recognisable groupings nor separate cleanly from one another, and the field emits no clear signal to policymakers about any of the actors holding them. Where preferences are structured into competing groupings, an actor's position is more clearly identifiable with that of the others sharing it, so that accommodating or overriding it means responding to more than the actor itself, and the political cost of doing so is correspondingly legible. Following this approach, the framework measures polarisation at the policy level, as a property of the consultation rather than of any participant, recording how cleanly the field of positions separates into internally cohesive and mutually opposed groupings, and falling as actors take positions that cut across any such division. Polarisation is therefore expected to moderate the effect of positional convergence.

H3 (Polarisation Amplification): The association between positional convergence and relative success is stronger the higher the level of polarisation within the consultation.

4. Research Design and Methodology

4.1 Data Collection and Processing

The dataset comprises 44 EU legislative files from the 2019–2024 legislature, generating 1,908 business stakeholder–policy observations from 1,839 unique business stakeholders. Files were drawn from the population on the Commission's Have Your Say portal by four criteria applied in sequence. The file must have been proposed and adopted within the 2019–2024 parliamentary term, so that the complete draft-to-final trajectory falls inside the observation window. Only regulations are retained, since a regulation is directly applicable and its adopted text is the operative legal instrument, whereas a directive is transposed into national law and its provisions cannot be traced to a single final text. The file must have attracted more than ten submissions from business actors, below which a network of positional relations cannot meaningfully be constructed. Finally, the regulation must be a self-standing body of new legal provisions, so amending regulations and recasts were excluded, because their provisions are intelligible only against a parent instrument. Files creating funding instruments or spending programmes, and files governing Union bodies or the coordination of national authorities, were also excluded, since they impose no requirements on business actors and so offer no provisions on which business preference attainment could be assessed.

The sample is restricted to business actors, meaning trade associations, companies, and corporate federations self-identifying as such in the Transparency Register or in their submissions. Besides, the sample only retains those submitting written position text in English; slides and presentations were excluded. The restriction to business is substantive rather than merely practical, since business actors dominate participation in EU consultations (Saurugger 2008; Schoenefeld 2021) and the literature this paper engages treats business largely as a single bloc, finding industry attainment to increase with the cohesion of the firms composing it (Chalmers 2020).

The resulting sample spans digital and internal market (18 policies), environment and climate (10), finance and anti-money-laundering (7), transport (2), health (2) and others, and the full range of polarisation from near-consensual (AI Act, $Q = 0.12$) to deeply divided (substances of human origin, $Q = 0.93$). Business stakeholders per policy range from 6 to 149 (median 26, mean 43.4), and 98.7% are matched to the EU Transparency Register (June 2024 snapshot).

Constructing provision-level measures requires linking three document sources: the draft regulation, the adopted regulation and each business stakeholder's submission, which differ in structure, granularity and register. Legal provisions are hierarchically numbered clauses; submissions are free-form position papers that may cite specific articles, address broad themes, or both. The procedure combines text extraction, semantic matching and large language model classification, and splits into two tracks. A *success track* builds the dependent variable, identifying how each provision changed between draft and final text and classifying each stakeholder's stance against that change; a *network track* builds the independent variable, comparing stakeholder positions pairwise. GPT-4.1-mini performs the segmentation of submissions into discrete policy claims; the full GPT-4.1 model performs the two classification tasks requiring contextual judgement, stance classification into five categories and dyadic position comparison into three. GPT-4-class models match or exceed expert human coders on political text classification, operating zero-shot (Heseltine & Clemm von Hohenberg 2024; Törnberg 2025), and their advantage is greatest on tasks requiring contextual inference of the kind involved in identifying positions relative to specific provisions. Appendix A specifies every stage, with parameter values, schemas and quality-assurance procedures. Appendix B sets out the evidence on validity, the treatment of model drift and correlated measurement error, and a human validation protocol in which two independent coders will classify stratified samples of the stance and dyadic classifications blind, with agreement to be reported as a Cohen's κ matrix covering human–human and human–model comparisons.

Two distinct observational questions are settled by this procedure, and the design answers each differently. Preferences are extracted from the stakeholder text at the level of the *paragraph* and matched to the provision of the legal text they address; the resulting provision–paragraph dyad is the unit of observation, at which a stance is classified and scored. Policy issues are identified from the legal text at the level of the *article*, which is the unit at which those provision-level stances are aggregated. The article is never observed directly, since no stance is classified against an article as such, and it enters only as the level at which the scores from its provisions are combined. This two-level treatment follows Bunea & Ibenskas (2017), whose unit of analysis is likewise a stakeholder paired with a discrete policy demand identified within submission text rather than with a legislative file as a whole. Where a stakeholder addresses a provision without citing it, the paragraph is matched to the legal text by semantic similarity, so which provision a position concerns is inferred in

those cases; the accuracy of that matching is assessed in the validation reported in Section 4.4.

4.2 Variable Construction

Dependent variable: relative success. Following Dür's (2008) taxonomy, lobbying outcomes can be measured in three ways: Process tracing reconstructs the pathway from demand to outcome but is confined to a few issues; attributed influence draws on self-reports or expert assessment, and so captures perceived rather than actual outcomes; and preference attainment, which this study adopts, compares stated positions against the policy output. Inputs are documentary, since the comparison is made against the two legal texts and the submission itself rather than against a participant's recollection or an expert's estimate, which makes the coding reproducible and the procedure scalable across 44 files. The measure records the correspondence between what an actor asked for and what the adopted text provides, and carries no implication that any particular submission produced the change.

The dependent variable (*dv_norm*) measures each business stakeholder's net attainment between the Commission proposal and the adopted regulation. For each pair linking an actor to a specific provision, GPT-4.1 classifies the business stakeholder's position relative to that text into one of five categories: aligned, where the stakeholder accepts the provision as drafted; divergent, where it accepts that the matter should be regulated but asks for a different rule; opposed, where it seeks the provision's removal; neutral or unclear, where the text engages the provision without taking a codeable position; and not the same issue, which removes false matches produced by the semantic matching. Establishing what happened to a provision requires the draft and adopted texts to be linked, and some draft provisions find no counterpart and have been dropped; others survive, whether unchanged or amended; and some provisions appear only in the adopted text and are new. Each provision's draft stance is then set against that outcome, and a scoring matrix assigns a value from -2 to $+2$ recording whether the outcome matched the stance, with ± 2 where a stance was decisively confirmed or reversed, as when an opposed stakeholder's target provision is removed or an aligned stakeholder's supported provision is dropped; ± 1 for partial shifts, where a provision moves between neutrality and a stance rather than reversing outright; and 0 where nothing changed or no outcome is attributable. Because the score records movement relative to the tabled text, a provision that survives unchanged returns no gain to the stakeholder that defended it and no loss to the one that contested it, so that an aligned stakeholder can only record a loss, where

the provision it supported is amended away, and a divergent stakeholder only a gain, where the amendment it sought is made. This classification, rather than a scaling method, is what allows the distinction argued for in Section 3.1 to be recorded, since divergence and opposition remain distinct categories. The matrix is given in full in Appendix A.

Constructing the variable requires the Commission proposal, the adopted regulation, and each business stakeholder submission to be brought together. The two legal texts are first divided into provisions, articles being subdivided into sections, subsections, sub-subsections and independent paragraphs. Each draft provision is then matched to its counterpart in the adopted regulation by embedding similarity above a fixed threshold. This hierarchical subdivision of the legal text allows the varying specificity of stakeholder preferences to be registered, from a whole article down to its subdivisions. The section is the broadest of these, and preferences too broad to attach below it are excluded at extraction. Aggregation is necessary because several legal provisions concern a single issue, and counting each separately would weight a stakeholder's outcome by how often it addresses it, since a demand may be pressed across several provisions of an article, or restated under different framings against the same provision, without the outcome differing. The same concern leads Bunea & Ibenskas (2017) to recode their finer-grained positions into a smaller set of demands, in their words "to account for the similarities among the positions". Appendix A.3 sets out the resolution rule and the scoring in full. The bounded article score makes the measure one of extensity, recording the number of issues on which a stakeholder prevailed rather than the degree to which it did so.

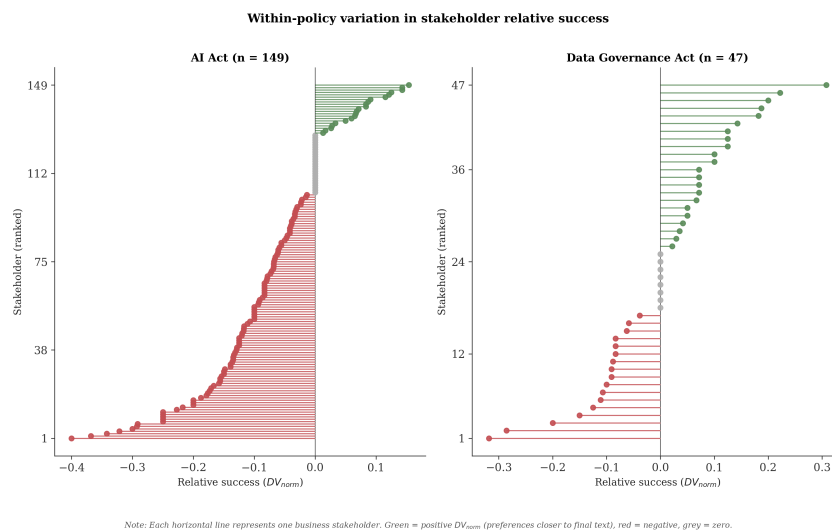


Figure 3. Within-policy variation in stakeholder relative success (DV_{norm}) for two contrasting regulations.

Note: Each horizontal line represents one business stakeholder, ordered by score. Green indicates positive outcomes (preferences closer to the final text than the draft), red indicates negative outcomes. The dispersion illustrates the substantial within-policy heterogeneity that the regression models seek to explain.

Independent variable: positional convergence. Business stakeholder pairs are formed where both actors commented on the same provision, identified through explicit citation or semantic matching. Their stances on that provision are compared and classified by GPT-4.1 into three dyadic relations, aligned, opposed or neutral, and the provision-level comparisons for a pair are summarised into a single normalised edge weight, $w_{\text{norm}} = (n_{\text{aligned}} - n_{\text{opposed}}) / (n_{\text{aligned}} + n_{\text{opposed}})$, positive where the two are net-aligned across their shared provisions and negative where net-opposed.

The choice of measure follows from the properties of the network under analysis. In this regard, the analysis follows Freeman (1978/79), who established that centrality is not a single quantity but "a family of measures" capturing structurally distinct properties; and Borgatti (2005), who shows that centrality measures embody "implicit models of how traffic flows" through a network, and that measures must accordingly be "matched to the kinds of flows that they are appropriate for" (p. 55). The first property bearing on that choice is that ties record correspondence between independently filed positions rather than any transfer between actors. This rules out flow-based measures such as betweenness and closeness, which describe control over, or proximity within, something that travels. The second is that those ties carry a valence, since a stakeholder may agree or disagree with those it engages. Conventional measures treat all ties as equivalent in valence and so cannot distinguish allies from opponents (Bonacich 1987), which is disabling where opposition counts against an actor; the signed-network centrality literature (Liu et al. 2020) instead treats agreement and opposition as contributions of opposite sign to the same quantity, which is the approach adopted here.

Degree and betweenness are nonetheless retained as predictors, precisely because they encode the assumptions this framework rejects and so allow it to be tested rather than asserted. Degree counts an actor's ties without regard to their sign, measuring how much of the field an actor engages rather than how far it agrees with it, and betweenness measures a brokerage that presupposes something travelling along the paths it counts. If the network operates as a structure of positional similarity rather than as a conduit, signed strength should predict success while betweenness should not, and the number of an actor's ties should confer no advantage of its own.

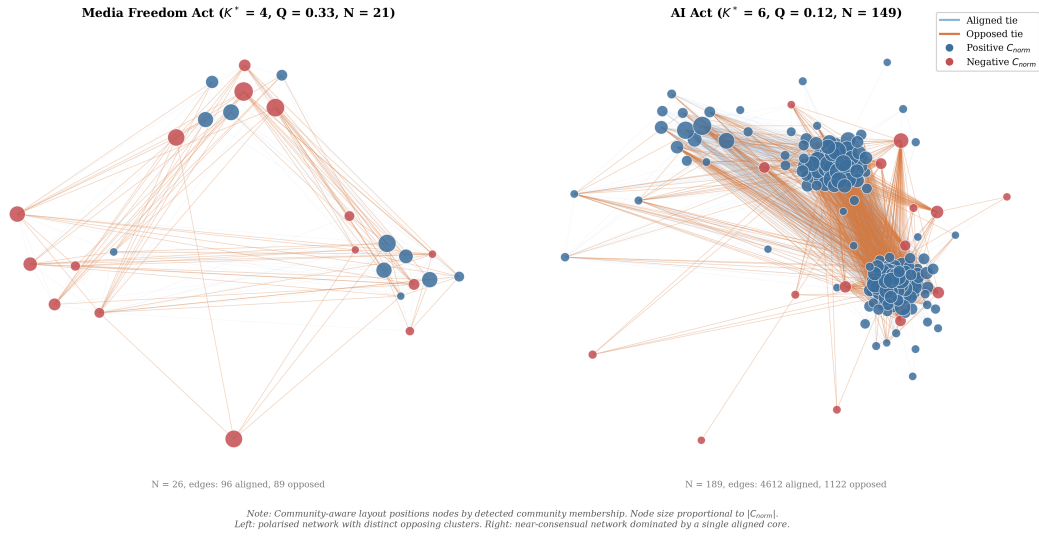


Figure 4. Signed network graphs illustrating contrasting network structures.

Note: Left: European Media Freedom Act ($Q = 0.33$, $N = 21$), a polarised network with clearly separated opposing clusters and a high proportion (57%) of net-opposed (red) stakeholders. Right: AI Act ($Q = 0.12$, $N = 149$), a near-consensual network dominated by a dense aligned core with few peripheral opponents. Blue = positive C_{norm} , red = negative C_{norm} , size proportional to $|C_{norm}|$.

Polarisation. This variable is measured by modularity Q , a scalar describing how cleanly a policy's signed network divides into internally cohesive, mutually opposed groups. A partition is obtained by applying the Louvain algorithm (Blondel et al. 2008) to each policy's aligned ties, without imposing a predetermined number of groups, and Q is then the signed modularity of that partition (Gómez et al. 2009), which rewards agreement within groups and opposition between them. Q enters the models as a context variable and moderator, continuously and, in the robustness checks, as a binary indicator above 0.40. It is not interpreted as a claim about how many sides a consultation contains or which side any actor belongs to. Values range from 0.12 (AI Act, near-consensual) to 0.93 (substances of human origin, deeply divided).

4.3 Controls

The first set of controls captures the organisational attributes most frequently linked to lobbying success in this literature, namely the resources an actor commands and the institutional access it enjoys (Bouwen 2002, 2004; Eising 2007; Dür et al. 2015). Resources are measured from the EU Transparency Register as members or full-time equivalent staff and annual lobbying expenditure, both log-transformed for extreme right skew. Access is proxied by Brussels office presence and accreditation to the European Parliament, with headquarters outside the EU indicating distance from the institutions.

A second set captures features of an actor's engagement. The stance profile records the proportions of a stakeholder's classified pairs falling into each stance category, with aligned as the reference. Umbrella membership indicates whether the actor is a trade association or federation rather than an individual company, testing whether formal aggregation of a broader constituency confers an advantage beyond what positional convergence captures. Submission intensity is the log word count of the submission, a crude proxy for the amount of information supplied (Klüver 2012).

A third set captures variation across policies rather than actors, since politicised files attract broader participation and more divided business stakeholder networks (Dür & Mateo 2012; Zürn et al. 2012), and the composition of the interest group population itself varies across domains, a pattern Coen & Katsaitis (2013) term "chameleon pluralism". Three dimensions associated with that variation are controlled. Consultation salience is the log number of feedback submissions a proposal received, proxying the breadth of attention to the file (Beyers & Arras 2020; Hanegraaff & Berkhout 2019), and ranges from 14 to 304 across the 44 policies. Policy sector is collapsed from EURLEX subject matter codes into eight categories with digital and internal market as reference, capturing the structural conditions under which politicisation operates, from domains of high public salience and organised opposition to those closer to expert-technocratic management (Bouwen 2002; Coen & Katsaitis 2013). Legislative complexity follows Senninger (2020) in distinguishing textual sophistication from structural scope, and is operationalised through the structural, linguistic and relational indicators of the EUPLEX framework (Hurka, Haag & Kaplaner 2021), five of which are defined in Appendix A.

Variable	Description	N	Mean	SD	Min	Max
DV_{norm}	Relative success (norm.)	1,966	-0.098	0.157	-1.000	1.000
C_{norm}	Signed coalition strength	1,966	10.211	14.986	-39.124	93.642
Degree	Network degree centrality	1,966	24.668	22.729	1.000	156.000
log(FTE)	Log members/FTE	1,966	1.135	0.834	0.000	3.792
log(Cost)	Log annual lobbying cost	1,966	10.028	5.007	0.000	21.330
EP accredited	EP accreditation passes	1,966	2.301	3.729	0.000	20.000
Brussels office	Has Brussels office (0/1)	1,966	0.211	0.408	0.000	1.000
Non-EU HQ	HQ outside EU (0/1)	1,966	0.130	0.336	0.000	1.000
Q	Signed modularity	1,966	0.330	0.151	0.119	0.932
log(Provisions)	Log N legal provisions	1,709	5.783	0.639	3.738	6.758
LIX score	Readability index	1,709	90.319	5.285	75.574	110.350
Participation ratio	Stakeholders / provisions	1,709	0.099	0.042	0.055	0.333

Table 2. Descriptive statistics for all variables used in the regression analysis.

Note: $N = 1,908$ for stakeholder-level variables; $N = 1,651$ across 36 policies for the complexity indicators. DV_{norm} is the normalised relative success score (range $[-1, 1]$); C_{norm} is the signed positional convergence measure derived from the network construction procedure.

4.4 Estimation Strategy

All models are linear mixed models estimated by maximum likelihood, with a random intercept by policy to account for the nesting of business stakeholders within legislative files; the ICC in the null model is 0.085, confirming that policy-level clustering is non-trivial. Continuous predictors are z-standardised, models are fitted in Python with statsmodels, and cases are complete on the variables entering each model, so N falls where Transparency Register covariates or complexity indicators are required. Models are built in blocks, beginning with network measures alone, then resource and access controls, then interactions with resources, polarisation, complexity and sector, and finally the stance profile, each compared to the previous by likelihood ratio test, so that the question is whether positional convergence survives the addition of controls, moderators and alternative explanations. Because the explanatory variable is derived from a network, the assumption that observations are independent conditional on the policy intercept justifies a test rather than an assertion, and a network autocorrelation specification (Doreian 1980; Leenders 2002) is reported in Section 5.4.

Both the dependent and the independent variable are derived from language-model classification, which raises questions of validity, of reliability, and of whether errors in the two could be correlated in ways that inflate the association between them.¹ Appendix B addresses these, setting out the evidence on the validity

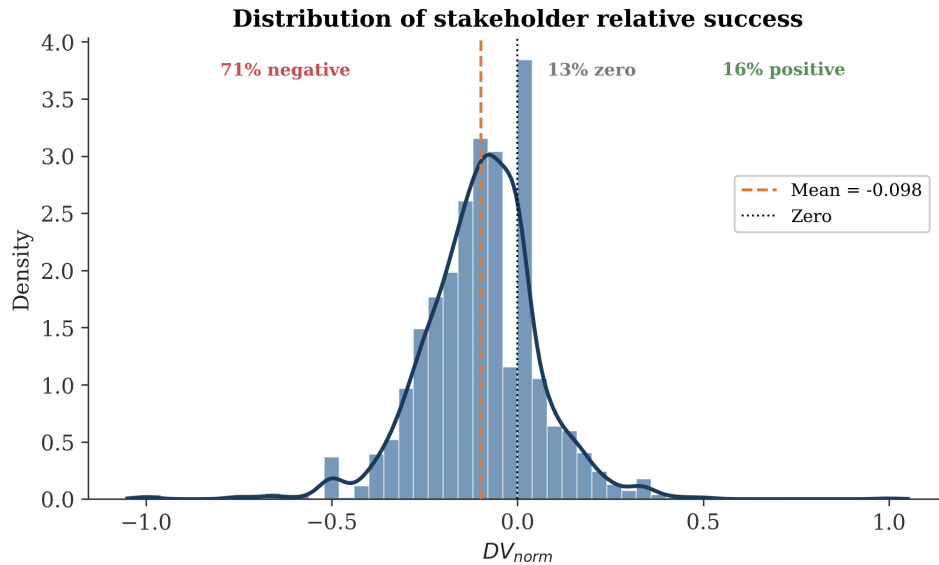
of language-model position estimation, the features of the present design that follow from it, the residual risk of correlated measurement error, and the human-validation protocol whose results will be reported as a Cohen's κ matrix.

5. Results

5.1 Descriptive Statistics

The analytical dataset comprises 1,908 business stakeholder–policy observations across 44 EU legislative files, covering 1,839 unique business stakeholders, of whom 97.1% appear in only one policy. Full descriptives are in [Table 2](#) and bivariate correlations in [Table 3](#).

Three features of the data matter for the analysis. First, the normalised relative success score averages -0.094 , with 71% of observations net negative, 13% exactly zero and 17% net positive; the location of the distribution reflects the scoring rules rather than the standing of business, since the measure records movement relative to the drafted text rather than an absolute level of attainment. That variation is also concentrated within policies rather than between them, ninety-one per cent of it sitting within ($ICC = 0.085$), which is what makes stakeholder-level explanation the appropriate object.



Note: $N = 1,966$ stakeholder-policy observations across 44 regulations. Kernel density estimate (solid line) fitted with Gaussian kernel.

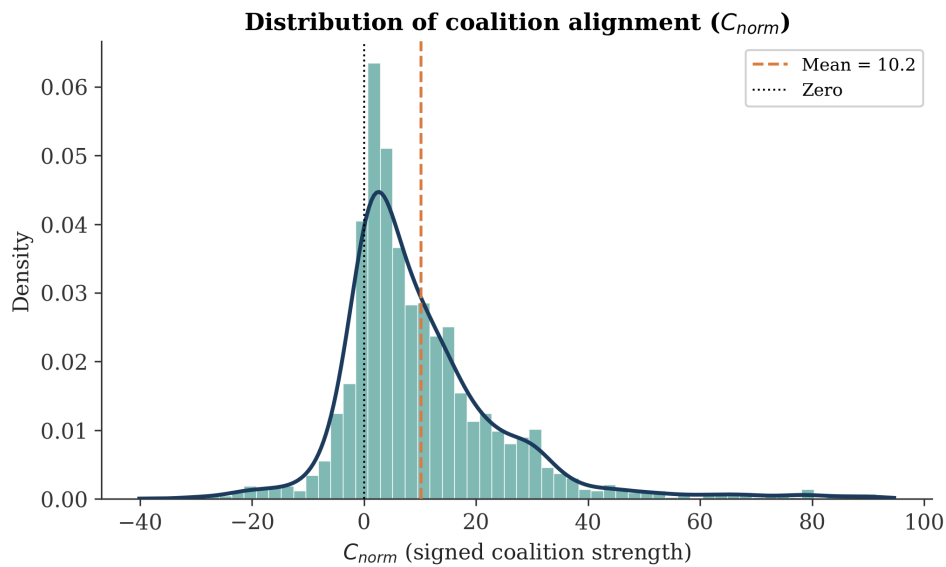
Figure 5. Distribution of business stakeholder relative success (DV_{norm}) across all 1,908 observations.

Note: The histogram shows a roughly symmetric unimodal distribution centred slightly below zero (mean = -0.094), indicating that most business stakeholders experience marginal losses between draft and final text. 71% of observations are negative, 13% are exactly zero, and 17% are positive.

Second, outright opposition is rare among the classified business stakeholder–provision pairs, of which 47.8% do not address the same issue, 32.6% are

divergent, 10.4% aligned, 8.6% neutral or unclear and only 0.6% opposed, so business stakeholders overwhelmingly propose alternative formulations rather than demanding deletion. The dyadic classifications show the same asymmetry, with 71.3% neutral, 20.4% aligned and 8.3% opposed.

Third, positional convergence is mechanically bound up with the structure of the consultation it is measured in. It averages 10.12 and is positive for 81.3% of observations, indicating more aligned than opposed ties on net, and it is highest in consensual policies and lowest in polarised ones, since where most dyads align every participant accumulates positive convergence. Policies span the full range of modularity Q , with 2 consensual ($Q < 0.20$), 19 moderate ($0.20 \leq Q < 0.40$) and 23 polarised ($Q \geq 0.40$). The interaction tested below turns on this, because convergence is most informative precisely where it is hardest to achieve.



Note: $C_{norm} = \sum_j W_{norm}(i, j)$ across all dyadic ties. Positive values indicate net-aligned stakeholders; negative values indicate net-opposed.

Figure 6. Distribution of signed positional convergence (C_{norm}) with kernel density overlay.

Note: The distribution is right-skewed with mean 10.1, indicating that most business stakeholders occupy net-positive positions within their consultation networks. Negative values (left of the dotted zero line) represent business stakeholders whose dyadic ties are predominantly opposed rather than aligned.

Resources show no bivariate association with success, consistent with the regression results reported below, though organisational size and institutional access correlate weakly with positional convergence ($r \leq 0.07$). Within this business-only sample, access is a minority characteristic, with 53.8% reporting no accredited parliamentary lobbyists and 78.8% no Brussels office, and lobbying expenditure is extremely right-skewed.

	DV _n	C _n	Deg	FTE	Cost	EP	Bru	nEU	Q	PR
DV _n	1.000									
C _n	0.036	1.000								
Deg	-0.013	0.712***	1.000							
FTE	0.023	0.085***	0.080***	1.000						
Cost	0.020	0.012	0.031	0.244***	1.000					
EP	0.012	0.072**	0.052*	0.706***	0.024	1.000				
Bru	0.015	0.060*	0.080***	0.398***	0.146***	0.180***	1.000			
nEU	-0.001	0.037	0.024	-0.045	0.044	-0.057*	0.187***	1.000		
Q	0.030	-0.290***	-0.419***	-0.021	-0.002	0.021	-0.035	-0.079**	1.000	
PR	0.081***	-0.080***	-0.070**	0.103***	0.031	0.113***	0.032	-0.040	0.274***	1.000

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

Table 3. Pairwise Pearson correlations among the main regression variables (lower triangle).

Note: Notable: *C_norm* and *Degree* are strongly correlated ($r = 0.71$, $p < 0.001$), motivating simultaneous inclusion in the regression models. *DV_norm* shows no bivariate association with resources but a small positive correlation with *C_norm* and participation ratio.

5.2 Main Results

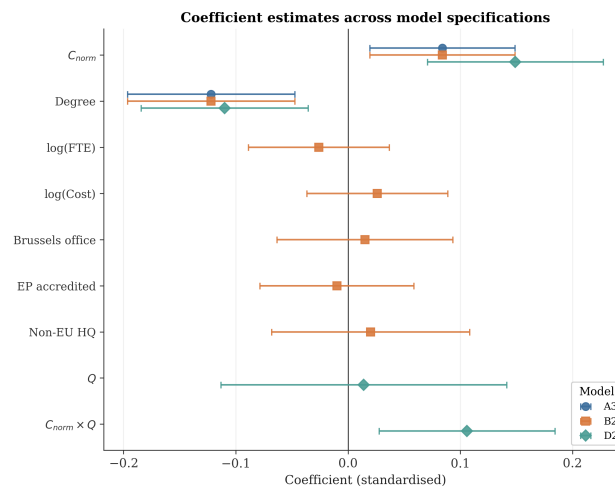
Positional convergence and degree. In the bivariate model (A1; Table 4), *C_norm* alone is non-significant. Once degree is controlled (A3), *C_norm* becomes significant ($\beta = +0.088$, $p < 0.01$) and degree shows a strong negative effect ($\beta = -0.129$, $p < 0.001$). The two measures are strongly correlated yet carry opposite signs, revealing a suppression effect. What matters is not how many actors a stakeholder is connected to but whether those connections form a coherent position within the field. Betweenness centrality is null across all specifications, indicating that brokerage position between disconnected groups confers no advantage in the consultation network.

In this regard, one standard deviation increase in positional convergence corresponds to a shift of roughly 0.013 in *dv_norm*, against a dependent variable whose standard deviation is 0.152, a real but modest effect. Comparing business stakeholders in the top and bottom quartiles of positional convergence gives a difference in mean success of 0.019. What the result establishes is that positional convergence predicts attainment where neither organisational resources nor the context of the file does. Adding resource and access controls (organisational size, lobbying expenditure, Brussels office, EP accreditation, non-EU headquarters) does not change the core finding (B2; Table 4). *C_norm* remains significant and all resource and access variables are individually non-significant, jointly adding no explanatory power over the baseline network model. The interaction between *C_norm*

and annual lobbying cost is marginally negative but does not reach conventional significance, hinting that resource-rich actors may benefit slightly less from positional convergence.

Sector fixed effects improve model fit (F6; Table 4), absorbing over half the between-policy variance and reducing the ICC from 0.086 to 0.042, but the C_norm effect is unchanged ($\beta = +0.088$, $p < 0.01$) and all resource controls remain null.

Polarisation interaction. Polarisation as a main effect is non-significant, meaning that the overall level of division within a consultation does not directly predict individual success. The $Q \times C_norm$ interaction, however, is significant ($\beta = +0.115$, $p < 0.01$), confirming that positional convergence matters more in polarised consultations where competing blocs crystallise around opposing positions. The binary polarisation interaction confirms this pattern ($\beta = +0.273$, $p < 0.01$), and in the interaction model the C_norm main effect increases to $\beta = +0.162$ ($p < 0.001$), suggesting that the baseline specification underestimates the centrality effect by averaging across consensual and polarised settings.



Note: Error bars show 95% confidence intervals. Coefficients standardised. A3 = baseline (C_norm + degree); B2 = + resource/access controls; D2 = + $Q \times C_norm$ interaction.

Figure 8. Forest plot of standardised coefficient estimates with 95% confidence intervals across three model specifications: A3 (baseline with C_norm and degree), B2 (adding resource and access controls), and D2 (adding polarisation interaction).

Note: The plot shows that the positive C_norm effect and the negative degree effect are stable across specifications, while resource and access controls cluster near zero.

Stance profile and divergence. The stance models are estimated on the observations with a draft-to-final transition ($N = 1,875$), and adding the stance profile substantially improves model fit. Divergence is the strongest predictor in any model estimated here ($\beta = -0.253$, $p < 0.001$), so that stakeholders whose positions diverge from the prevailing regulatory orientation face a substantial penalty in what they

obtain. Opposition does not reach significance, which is unsurprising given its rarity, since only 8.7% of business stakeholders have any opposition pairs at all. Neutral stance is modestly positive ($\beta = +0.058$, $p < 0.01$), suggesting that business stakeholders who engage a provision without taking a clear position on it fare slightly better than those who contest it. The $C_norm \times$ divergence interaction is positive but falls short of conventional significance ($\beta = +0.049$, $p = 0.061$), offering only suggestive support for the expectation that positional convergence compensates for a divergent stance.

	(1)	(2)	(3)
	<i>Baseline</i>	<i>Polarisation</i>	<i>Stance</i>
C_norm	+0.084** (0.032)	+0.149*** (0.039)	+0.074* (0.031)
Degree	-0.122*** (0.036)	-0.102** (0.037)	-0.069. (0.035)
<i>Controls</i>			
log(FTE)		-0.033 (0.030)	-0.021 (0.029)
log(Cost)		+0.019 (0.023)	+0.009 (0.022)
Brussels office		+0.014 (0.060)	+0.001 (0.058)
EP accredited		+0.049 (0.059)	+0.030 (0.057)
Non-EU HQ		-0.011 (0.066)	+0.021 (0.065)
<i>Polarisation</i>			
Q (signed)		+0.072 (0.049)	
C_norm $\times$ Q		+0.106** (0.036)	
<i>Stance mechanism</i>			
Stance: divergence			-0.257*** (0.023)
Stance: opposition			+0.021 (0.021)
Stance: neutral			+0.058** (0.022)
C_norm $\times$ Div.			+0.056* (0.025)
N	1,966	1,966	1,932
ICC	0.089	0.090	0.064
AIC	5,477	5,480	5,233

Standard errors in parentheses. . $p < 0.10$, * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$.
Random intercept by policy. All models include resource/access controls where shown.
Model 3 uses draft-transition subsample ($N = 1,932$). Full model progression in Table A4.

Table 4. Regression results.

Note: Linear mixed models with random intercept by policy. A1: bivariate (C_norm only). A3: adds network degree, revealing the suppression effect. B2: adds resource and access controls. F6: adds sector fixed effects. D2: polarisation interaction ($Q \times C_norm$). G3: stance profile and $C_norm \times$ divergence interaction, estimated on the 1,875 observations with a draft-to-final transition. Standard errors in parentheses. Extended models (complexity, side-level, robustness) in Table A4.

5.3 Politicisation

Section 4.3 introduced three policy-level dimensions along which consultations vary, and they behave differently. Legislative complexity does not predict success on any of its five measures, and the strongest complexity interaction, between the participation ratio and positional convergence, falls short of conventional significance. Consultation salience is likewise null as a main effect, though its interaction with positional convergence is negative ($\beta = -0.071$, $p < 0.05$), which would suggest that convergence is somewhat less effective where a wider audience is engaged, a pattern consistent with the quiet-politics expectation that concentrated interests do better out of public view, but too weakly evidenced here to be read as support for it.

Policy sector is the exception. Sector fixed effects absorb over half the between-policy variance, reducing the ICC from 0.086 to 0.042, and environment and climate files carry a strong negative intercept ($\beta = -0.305$, $p < 0.01$), so that business stakeholders attain systematically less in this domain than elsewhere. Finance and anti-money-laundering files show a more favourable but non-significant pattern. What the result does establish is that the positional-convergence effect is not an artefact of sector composition, being unchanged in magnitude and significance once sector is absorbed.

5.4 Further Tests

Three additional analyses address specific concerns raised in the literature ([Table A4](#)).

Submission intensity. Text length of the consultation response is a proxy for the amount of information supplied to policymakers (Klüver 2012). The main effect on success is null, and the interaction with C_norm is negative ($\beta = -0.001$, $p = 0.001$), so that the returns to positional convergence fall as submissions lengthen. The volume of text supplied therefore does not substitute for positional convergence, and a long submission poorly embedded within the broader demand structure appears to dilute rather than reinforce it. The measure is nonetheless a crude one, since length captures how much a stakeholder wrote, not how informative what it wrote was.

Umbrella membership. Formal umbrella organisation membership (trade associations, EU-level federations) has no main effect on success and no interaction with C_norm . Umbrella organisations are not associated with higher centrality either, indicating that formal organisational aggregation does not confer an advantage beyond what positional convergence already captures.

Interdependence among networked observations. The models estimated above treat business stakeholders as independent once policy-level clustering is absorbed by

the random intercept, which is a strong assumption in a design whose explanatory variable is relational, since two actors aligned on the same provisions may plausibly have correlated outcomes for reasons the specification does not capture. This is tested directly with a network autocorrelation model (Doreian 1980; Leenders 2002), which adds an autoregressive term in the outcomes of network neighbours, $y = \rho W y + X\beta + \epsilon$, where W is the signed adjacency matrix, block-diagonal by policy so that no ties cross legislative files, and row-normalised. The autocorrelation parameter is substantial and highly significant ($\rho = 0.352$, $p < 0.001$), confirming that outcomes are not independent across the network. The positional convergence coefficient nonetheless survives, and in fact strengthens, under this specification ($\beta = +0.105$, $z = 3.31$), so the result is not an artefact of unmodelled dependence among connected business stakeholders. Estimated separately on aligned and opposing ties, interdependence operates almost entirely through convergence ($\rho = 0.363$, $p < 0.001$) rather than through opposition ($\rho = 0.083$, $p = 0.017$).

6. Discussion

6.1 Positional Convergence, Network Structure, and Divergence

The main finding is that where a business stakeholder sits in the field of stated positions predicts whether its preferences reach the final text (H1), while neither its organisational attributes nor the context of the file predicts it directly. This is consistent with the signalling account, on which what a consultation indicates to policymakers is the breadth and depth of the support a proposal commands (Kingdon 1995; Kollman 1998), a property of the distribution of stated positions rather than of any actor's own advocacy. While the coalition literature locates this property in formal association or deliberate cooperation (Mahoney 2007b; Junk 2019), part of it looks as well at the consensus across the messages members send (Nelson & Yackee 2012) and at which actors hold the same view (Phinney 2017). The results support this second reading, as positional convergence is computed from independently filed submissions, with no formal tie between the actors it links, so coincident positions register whether or not organisational cooperation underlies them (Baumgartner et al. 2009; Atikcan & Chalmers 2019), each actor being characterised by its own location in that structure rather than by the properties of a grouping (Bunea 2013).

A tie records a correspondence between two independently stated positions rather than a transfer between the actors holding them, so the network structure is one of similarity rather than of flow (Borgatti 2005). Betweenness, which identifies actors able to broker between others, is accordingly null. Stakeholders file with the institutions rather than negotiate with one another, so bridging otherwise separate parts of the field yields no advantage. Degree, the number of actors whose demands meet an actor's regardless of the agreement or disagreement, is negatively associated with attainment. Engaging the provisions many others engage confers nothing in itself, since what an actor gains depends on whether those others support its demands or contest them. Whether an actor succeeds is therefore explained by how coherently its demands sit within those of the others addressing the same provisions, not by how many of them it stands among.

The polarisation interaction (H3) extends this reading from the demand to the field it is filed into, since whether the support a demand attracts registers at all depends on how the rest of the field is arranged. Where preferences are structured into competing groupings the field emits a legible signal and positional convergence differentiates outcomes; where it is weakly divided, convergence is high across most actors and conveys little about any of them. This is why positional convergence and

polarisation are treated as measuring different things, the first being a property of an actor, the second of the setting that determines what the first implies.

This study addresses the legitimisation stage, which the consultation literature has left largely unobserved. The characteristics of this stage motivated two departures from established practices, as issues are inferred from stakeholder engagement rather than taken from the document by which the institution organised the consultation (Chalmers 2020; Bunea, Wüest & Lipcean 2025); and preferences are recorded by their substantive content rather than by their position on a scale of regulatory stringency (Bunea 2013; Bernhagen, Dür & Marshall 2014; Dür, Bernhagen & Marshall 2015). This second decision allows support for a provision, divergence from it and opposition to it to be distinguished, and the results show that supportive and divergent demands attain at systematically different rates. Business stakeholders whose demands diverge from the prevailing regulatory orientation attain significantly less than those supporting the text as drafted (H2a), and the asymmetry is directly observable, since a provision a stakeholder supported survives in the form it supported in 31 per cent of cases, whereas a provision it contested moves to the position it sought in 4 per cent. This accords with Bunea (2013), who reports that status-quo demands attain more than demands for change, a pattern the literature attributes to the general advantage defenders of an existing position enjoy (Baumgartner et al. 2009). Positional convergence appears to matter more for divergent stakeholders (H2b), though the interaction falls short of conventional significance. This is consistent with the same reading, since support from the field is what a demand needs when it requires the text to move, and is of little consequence when it does not.

Neither the resources an actor commands nor the access it enjoys predicts attainment in the study. Organisational size, annual lobbying expenditure, a Brussels office, accreditation to the Parliament and headquarters outside the Union are individually and jointly non-significant, as is the length of the submission filed. In this regard, three studies measuring the same organisational properties report similar results. Mahoney (2007a) finds resources unrelated to preference attainment once issue context is controlled; Dür, Bernhagen and Marshall (2015) attribute success in adopted legislation to group type and issue conflict rather than organisational endowment; and Chalmers (2020) reports firm resources correlating negatively with success. What the present results add is the stage and the level of measurement, since none of these observes legitimisation, and each records outcomes for the file rather than issue by issue. It does not follow that the information a stakeholder supplies is

irrelevant to lobbying success, since only the volume of that information is observed here. Bunea et al. (2026) show that the informational quality of public comments varies along dimensions only weakly related to volume, so that the amount supplied and the quality of what is supplied are distinct quantities.

The first limitation of this study concerns scope, since the framework applies to the observable field of stated positions and preferences communicated through channels that leave no public record are not captured. The analysis is further restricted to business actors and to regulations, as extending the design to non-business participants would allow the business–civil society cleavage to be examined at provision level, and would test whether the relational mechanism identified here operates across a more heterogeneous field. The third of these scope limitations follows from the measure, since directives are not covered and the sample covers only the instruments through which the Union legislates directly and the findings may not extend to files where implementation is left to the member states. Although organisational resources are the standard measure in this literature, as the Transparency Register reports annual expenditure for the organisation as a whole rather than for any single file, policy-specific financial investment remains unobserved. Consequently, while the resource nulls reported here are robust across specifications, a measure defined at the level of the organisation rather than the file may not capture what a group commits to a particular proposal. The politicisation results are likewise inconclusive, since legislative complexity is unrelated to success on all five of its measures, consultation salience is null as a main effect, and only policy sector discriminates, environment and climate files carrying a negative intercept. These three dimensions are properties of the policy, of which there are forty-four, so an absent association may reflect the number of files observed as much as the absence of an effect.

7. Conclusion

Among business stakeholders participating in EU public consultations, whose preferences prevail is predicted by where an actor's demands sit within the structure of positions the consultation produces. Positional convergence predicts success in every specification that holds constant how many actors a stakeholder engages (H1), while the resources an organisation commands, the access it enjoys, its umbrella membership and the length of the submission it files do not, individually or jointly, and neither the complexity of a file nor its salience predicts attainment directly. Business stakeholders whose demands diverge from the prevailing regulatory orientation attain significantly less than those supporting the text as drafted (H2a), and the convergence effect is amplified where the field divides into mutually opposed groupings (H3), while the expectation that convergence matters more for divergent actors falls short of conventional significance (H2b). Three contributions follow. The first is a measure, since positional convergence records where an actor stands in the field of stated positions, counting agreement in its favour and opposition against it, and so extends the binary indicator of median position Bunea (2013) established at formulation, while differing from measures defined for the grouping as a whole (Klüver 2012, 2013), which assign every member the same value. The alternatives fail for reasons the framework anticipates, since brokerage confers nothing where stakeholders file with the institutions rather than negotiate with one another, and engaging the same provisions as many others carries a penalty rather than an advantage. The second is the stage, since carrying the account to legitimation required issues to be inferred from stakeholder engagement, an article becoming an issue only when stakeholders address it, and demands to be recorded by what they ask for rather than by their place on a scale of regulatory stringency, which is what allows support, divergence and opposition to be told apart. The third is substantive, since the divergence penalty follows from what is distinctive about legitimation, the text under comment being a position the Commission has just adopted and will defend through the negotiation that follows, so that business influence is exercised most reliably in defence of a text rather than against it. Outcomes are moreover interdependent across the network almost entirely through agreement rather than through conflict, so that business is structured less as a set of contending camps than as a field in which shared positions carry actors together.

The convergence effect is modest, and most of the variance in who obtains what remains unexplained, which is consistent with a process in which the fate of a given provision turns on considerations no consultation-based measure observes, so that

positional convergence predicts attainment where resources measured at the level of the organisation do not, without determining it. Regarding other potential improvements, extending the analysis to non-business participants would test whether the mechanism operates across a more heterogeneous field. Besides, data on what a group commits to a particular proposal, rather than to its activities as a whole, would establish whether the resource nulls reported here reflect an absence of resource effects or the imprecision of organisation-level measures. Measuring the informational quality of submissions rather than their volume would test what the present null on length leaves open (Bunea et al. 2026). Combining consultation texts with data on meetings, amendment authorship and committee proceedings would separate the pathways by which a stated position comes to be reflected in the adopted text.

Notes

1. Language-model estimates of political positions have been validated against expert-survey benchmarks at correlations of .87 to .92, against a human-coder benchmark of .3 to .5 for the procedures they replace (Benoit et al. 2026; Mikhaylov et al. 2012). Classification errors do not bias regression estimates unless systematically correlated with the variables of interest (Heseltine & Clemm von Hohenberg 2024); here the two tasks classify structurally different constructs from different inputs, which limits though does not eliminate that risk (Abdurahman et al. 2025; McLaren et al. 2026).

8. References

- Abdurahman, S., et al. (2025). Best practices for text annotation with large language models. *arXiv preprint*.
- Atikcan, E.Ö., & Chalmers, A.W. (2019). Choosing lobbying sides: The General Data Protection Regulation of the European Union. *Journal of Public Policy*, 39(4), 543–564.
- Baumgartner, F.R., Berry, J.M., Hojnacki, M., Kimball, D.C., & Leech, B.L. (2009). *Lobbying and Policy Change: Who Wins, Who Loses, and Why*. University of Chicago Press.
- Bell, S., & Hindmoor, A. (2009). *Rethinking Governance: The Centrality of the State in Modern Society*. Cambridge University Press.
- Benoit, K., De Marchi, S., Laver, C., Laver, M., & Ma, J. (2026). Using large language models to analyze political texts through natural language understanding. *American Journal of Political Science*. Advance online publication. <https://doi.org/10.1111/ajps.70050>
- Beyers, J., & Arras, S. (2020). Stakeholder lobbying and policy outcomes: The access–influence distinction. *Journal of European Public Policy*, 27(8), 1143–1162.
- Beyers, J., & De Bruycker, I. (2018). Lobbying makes (strange) bedfellows: Explaining the formation and composition of lobbying coalitions in EU legislative politics. *Political Studies*, 66(4), 959–984.
- Bernhagen, P., Dür, A., & Marshall, D. (2014). Measuring lobbying success spatially. *Interest Groups & Advocacy*, 3(2), 202–218.
- Binderkrantz, A.S., & Pedersen, H.H. (2019). The lobbying success of citizen groups and business interests in Denmark. *Acta Politica*, 54, 583–599.
- Björnsson, C.H. (1968). *Läsbarhet*. Stockholm: Liber.
- Blondel, V.D., Guillaume, J.-L., Lambiotte, R., & Lefebvre, E. (2008). Fast unfolding of communities in large networks. *Journal of Statistical Mechanics: Theory and Experiment*, 2008(10), P10008.
- Bonacich, P. (1987). Power and centrality: A family of measures. *American Journal of Sociology*, 92(5), 1170–1182.
- Borgatti, S.P. (2005). Centrality and network flow. *Social Networks*, 27(1), 55–71.
- Bouwen, P. (2002). Corporate lobbying in the EU: The logic of access. *Journal of European Public Policy*, 9(3), 365–390.
- Bouwen, P. (2004). Exchanging access goods for access: Business lobbying in EU institutions. *European Journal of Political Research*, 43(3), 337–369.
- Bunea, A. (2013). Issues, preferences and ties: Determinants of interest groups' preference attainment in the EU environmental policy. *Journal of European Public Policy*, 20(4), 552–570.

- Bunea, A., & Ibenskas, R. (2015). Quantitative text analysis and the study of EU lobbying and interest groups. *European Union Politics*, 16(3), 429–455.
- Bunea, A., & Ibenskas, R. (2017). Unveiling patterns of contestation over Better Regulation reforms in the European Union. *Public Administration*, 95(3), 589–604.
- Bunea, A., Wüest, R., & Lipcean, S. (2025). Mapping the policy space of public consultations: Evidence from the European Union. *Journal of European Public Policy*, 32(3), 755–783.
- Bunea, A., & Lipcean, S. (2026). Does public participation in supranational policymaking enhance the European Commission's legislative agenda-setting success? *Journal of European Public Policy*. Advance online publication. <https://doi.org/10.1080/13501763.2026.2623923>
- Bunea, A., Lipcean, S., & Rauh, C. (2026). Responsive to what? Explaining the information quality of public comments on bureaucratic policymaking using a text-as-data approach. *Governance*, 39(2), e70122.
- Chalmers, A.W. (2020). Unity and conflict: Explaining financial industry lobbying success in European Union public consultations. *Regulation & Governance*, 14(3), 391–408.
- Coen, D., & Katsaitis, A. (2013). Chameleon pluralism in the EU: An empirical study of the European Commission interest group density and diversity across policy domains. *Journal of European Public Policy*, 20(8), 1104–1119.
- De Bruycker, I., & Beyers, J. (2019). Lobbying strategies and success: Inside and outside lobbying in EU legislative politics. *European Political Science Review*, 11(1), 57–74.
- Downs, A. (1957). *An Economic Theory of Democracy*. New York: Harper & Row.
- Doreian, P. (1980). Linear models with spatially distributed data: Spatial disturbances or spatial effects? *Sociological Methods & Research*, 9(1), 29–60.
- Dür, A. (2008). Interest groups in the European Union: How powerful are they? *West European Politics*, 31(6), 1212–1230.
- Dür, A., & Mateo, G. (2012). Who lobbies the EU? National interest groups in a multilevel polity. *Journal of European Public Policy*, 19(7), 969–987.
- Dür, A., & Mateo, G. (2016). *Insiders versus Outsiders: Interest Group Politics in Multilevel Europe*. Oxford University Press.
- Dür, A., Bernhagen, P., & Marshall, D. (2015). Interest group success in the European Union: When (and why) does business lose? *Comparative Political Studies*, 48(8), 951–983.
- Eising, R. (2007). Institutional context, organisational resources and strategic choices: Interest group access in the EU. *European Union Politics*, 8(3), 329–362.
- European Commission (2017). *Better Regulation Toolbox*. Commission Staff Working Document, SWD(2017) 250.
- European Commission (2021). *Better Regulation Guidelines*. Commission Staff Working Document, SWD(2021) 305 final, 3 November 2021.
- European Court of Auditors (2019). *'Have your say!': Commission's public consultations engage citizens, but fall short of outreach activities*. Special Report No 14/2019.

- Freeman, L.C. (1978/79). Centrality in social networks: Conceptual clarification. *Social Networks*, 1(3), 215–239.
- Gómez, S., Jensen, P., & Arenas, A. (2009). Analysis of community structure in networks of correlated data. *Physical Review E*, 80(1), 016114.
- Halterman, A., & Keith, K. (2026). Codebook LLMs: Evaluating LLMs as measurement tools for political science. *Political Analysis*.
- Hanegraaff, M., & Berkhout, J. (2019). More business as usual? Business bias across issues and institutions in the EU. *Journal of European Public Policy*, 27(8), 1143–1162.
- Heseltine, M., & Clemm von Hohenberg, B. (2024). Large language models as a substitute for human experts in annotating political text. *Research & Politics*, 11(1).
- Hurka, S., Haag, M., & Kaplaner, C. (2021). Policy complexity in the European Union, 1993–2020. *Journal of European Public Policy*, 29(11), 1747–1768.
- Hula, K. (1999). *Lobbying Together: Interest Group Coalitions in Legislative Politics*. Georgetown University Press.
- Junk, W.M. (2019). When diversity works: The effects of coalition composition on the success of lobbying coalitions. *American Journal of Political Science*, 63(3), 660–674.
- Kingdon, J.W. (1995). *Agendas, alternatives, and public policies* (2nd ed.). New York: HarperCollins.
- Klüver, H. (2009). Measuring interest group influence using quantitative text analysis. *European Union Politics*, 10(4), 535–549.
- Klüver, H. (2011). The contextual nature of lobbying: Explaining lobbying success in the European Union. *European Union Politics*, 12(4), 483–506.
- Klüver, H. (2012). Lobbying in coalitions: Interest group influence on EU policy-making. *European Journal of Political Research*, 51(4), 424–445.
- Klüver, H. (2013). *Lobbying in the European Union: Interest Groups, Lobbying Coalitions, and Policy Change*. Oxford University Press.
- Leenders, R.T.A.J. (2002). Modeling social influence through network autocorrelation: Constructing the weight matrix. *Social Networks*, 24(1), 21–47.
- Lin, Y., & Zhang, H. (2025). From human coding to LLM coding: A framework for social science text analysis. *Computational Communication Research*.
- Kollman, K. (1998). *Outside lobbying: Public opinion and interest group strategies*. Princeton University Press.
- Liu, N.F., Lin, K., Hewitt, J., Paranjape, A., Bevilacqua, M., Petroni, F., & Liang, P. (2023). Lost in the middle: How language models use long contexts. *Transactions of the Association for Computational Linguistics*, 12, 157–173.
- Liu, Z., Ma, J., Zeng, Y., Yang, L., Huang, Q., & Wu, H. (2020). On the centrality of actors in networks with signed relations. *Information Sciences*, 512, 613–630.
- Mahoney, C. (2007a). Lobbying success in the United States and the EU. *Journal of Public Policy*, 27(1), 35–56.

- Mahoney, C. (2007b). Networking vs. allying: The decision of interest groups to join coalitions in the US and the EU. *Journal of European Public Policy*, 14(3), 366–383.
- McLaren, J., et al. (2026). Researcher degrees of freedom in LLM-assisted content analysis. *Political Communication*.
- Mikhaylov, S., Laver, M., & Benoit, K. (2012). Coder reliability and misclassification in the human coding of party manifestos. *Political Analysis*, 20(1), 78–91.
- Nelson, D., & Yackee, S.W. (2012). Lobbying coalitions and government policy change: An analysis of federal agency rulemaking. *Journal of Politics*, 74(2), 339–353.
- Phinney, R. (2017). *Strange Bedfellows: Interest Group Coalitions, Diverse Partners, and Influence in American Social Policy*. Cambridge University Press.
- Roy, O., & Vetterli, M. (2007). The effective rank: A measure of effective dimensionality. *15th European Signal Processing Conference*, 606–610.
- Saurugger, S. (2008). Interest groups and democracy in the EU. *West European Politics*, 31(6), 1274–1291.
- Schoenefeld, J.J. (2021). Interest groups, NGOs, and the political economy of climate policy. *Wiley Interdisciplinary Reviews: Climate Change*, 12(6).
- Senninger, R. (2020). What makes policy complex? Examining the role of information in legislative politics. *European Political Science Review*, 12(4), 461–479.
- Thomson, R. (2009). Actor alignments in the European Union before and after enlargement. *European Journal of Political Research*, 48(6), 756–781.
- Törnberg, P. (2025). ChatGPT-4 outperforms experts and crowd workers in annotating political Twitter messages with zero-shot learning. *arXiv preprint*.
- Woll, C. (2007). Leading the dance? Power and political resources of business lobbyists. *Journal of Public Policy*, 27(1), 57–78.
- Ziems, C., Held, W., Shaikh, O., Chen, J., Zhang, Z., & Yang, D. (2024). Can large language models transform computational social science? *Computational Linguistics*, 50(1), 237–291.
- Zürn, M., Binder, M., & Ecker-Ehrhardt, M. (2012). International authority and its politicization. *International Theory*, 4(1), 69–106.

Appendix A. Technical Pipeline Specification: DV and IV Construction

A.1 Legal Text Extraction and Alignment (Step 1)

Extraction. The Commission proposal (draft) and adopted regulation (final) are processed from PDF into structured CSV files. The extraction scripts (EC_cleaning_splitting.py for drafts, Final_reg_cleaning_splitting.py for finals) use pdfplumber's word-level extraction with custom line reconstruction that groups words by vertical position and detects paragraph boundaries through inter-line spacing gaps. The raw text is then cleaned through a series of regex operations: institutional headers and footers (EN EN markers, standalone page numbers) are removed, inline structural leakage (CHAPTER/SECTION/TITLE headings) is stripped, footnotes and OJ references are deleted, and financial statements and annexes are truncated. After cleaning, articles are identified by detecting "Article N" heading patterns and extracting any associated title text.

Each article is then hierarchically decomposed into its constituent provisions. Numbered paragraphs marked with (1), (2), (3) become *sections*; lettered lists (a), (b), (c) within sections become *subsections*; and roman-numeral enumerations (i), (ii), (iii) within lettered items become *subsubsections*. Standalone paragraphs within a section that carry no letter or numeral marker are labelled p1, p2, etc. For each section, the script also emits a `full` entry containing the complete section text before decomposition. Articles are not analysed as single units because they are too large; instead, each article's content is fully covered through its finer subdivisions. Each provision is assigned a canonical 4-level `unit_id` following the format `{version}:{article}.{section}.{subsection}.{subsubsection}`: for instance, `draft:2.2.a.ii` isolates subsubsection (ii) of subsection (a) of section 2 of Article 2, while `draft:3.none.p1.none` captures the first paragraph of Article 3 when it has no numbered sections. Several edge cases require specialised handling: nested numeric enumerations inside lettered items are protected from being misidentified as section boundaries, page-break soft wraps (where a sentence continues on the next page after a blank line) are detected and merged back together, and cross-instrument legal references (e.g., citations to TFEU articles) are distinguished from the regulation's own structural markers to avoid false splits. Definitions articles receive separate treatment, mapping each lettered definition to a numbered section. Where a section's `full` entry and its first unlettered paragraph

carry identical text, the paragraph is dropped, so that the same content is not recorded twice.

Alignment. Draft and final provisions are aligned to establish legal continuity. Using sentence-transformers (all-MiniLM-L6-v2), each provision is embedded into a 384-dimensional vector. A blocked computation first identifies the top 50 candidate matches for each provision on both sides (draft→final and final→draft), processing embeddings in chunks to manage memory. A two-wave greedy unique assignment then selects best-match pairs from these candidates: all candidate edges are sorted by cosine similarity in descending order, and each source and target can only be assigned once, ensuring one-to-one matching. Wave 1 assigns each draft provision to its best available final provision; Wave 2 independently assigns each final provision to its best available draft provision. Pairs with cosine similarity below $\tau = 0.70$ (TAU_CONT_LEGAL) are treated as unmatched. Duplicate continuity pairs across waves are deduplicated, keeping the draft-side row as canonical. A consistency fix relabels any final-side continuity row whose matched draft provision was classified as "dropped" in Wave 1, since these arise from asymmetric threshold crossings between the two independent waves. Each provision receives a continuity label: *continuity* (matched above threshold), *dropped* (draft with no final match), or *new_final* (final with no draft counterpart). The continuity map is the backbone of the DV, since it determines which draft provisions have trackable outcomes.

A.2 Stakeholder Text Processing (Step 2)

Cleaning and chunking. Raw stakeholder PDFs are processed into cleaned paragraph-level JSON chunks. PDFs are read with PyMuPDF at the block level, with blocks sorted into reading order (including two-column layout detection for documents with side-by-side text). Repeated headers and footers are identified by normalising the top and bottom lines of each page and removing any line that appears on at least half of all pages. The text then passes through several cleaning steps: standalone page numbers and institutional markers are removed, footnote markers and OJ references are stripped, hyphenated line breaks are rejoined, and whitespace is normalised. A line-repair pass merges broken lines caused by PDF extraction artifacts: isolated single-word fragments, lines ending with connective words ("and," "of," "the"), enumeration stubs, and cross-page sentence splits are detected through heuristic rules and joined to adjacent lines. The repaired text is then split into paragraphs using blank-line boundaries, with list items (lettered, numbered, bulleted) kept attached to their parent sentence. A multi-pass paragraph repair stage merges broken paragraph pairs (incomplete sentences followed by continuation lines), coalesces multi-line titles, and separates heading–body fusions where a section heading was concatenated with the body text that follows it. The resulting paragraphs are grouped into token-bounded chunks targeting 9,000 tokens with a hard cap at 11,000 tokens and 1-paragraph overlap between chunks, using tiktoken (cl100k_base) for token estimation. Each chunk is stored as a separate JSON file. A stakeholder registry CSV is also produced, mapping each PDF to a canonical stakeholder ID derived from either the filename or, when the filename is a hash, from heuristic extraction of the organisation name in the document's opening lines.

Segmentation and filtering. GPT-4.1-mini (temperature = 0) with OpenAI's Structured Outputs (strict JSON schema) segments and filters each chunk into atomic, politically substantive stakeholder units. The model extracts concrete amendment requests, support or opposition directed at identifiable provisions or policy elements, requests for clarification, deletion, addition, narrowing, or broadening, and provision-anchorable interpretations of how the regulation should operate. It discards material that does not meet these criteria: general support or opposition directed at the overall initiative without a provision-level anchor, purely descriptive or contextual statements that characterise the regulation without taking a position, rhetorical framing, administrative boilerplate, and generic policy claims that cannot be tied to a concrete element of the regulation. Each extracted unit is assigned an ID, text content, a binary explicit-reference flag indicating whether the stakeholder cites a specific

provision of the target regulation and, when the flag is set, the cited legal location (article/section/subsection/subsubsection), all constrained by a strict schema. When the stakeholder cites a provision from another legal instrument (e.g., TFEU, another directive), the legal-location fields are left empty, being filled only with references to the target regulation under consultation. When a single substantive claim is directed at multiple articles, each cited article generates a separate unit.

A.3 DV Track: Stakeholder–Legal Matching, Stance Classification, and Scoring (Steps 3–4)

Matching (Step 3). Each stakeholder unit is matched to legal provisions through multiple sources, operating in a single pass that produces both draft and final matches for each unit. First, all stakeholder units, draft provisions, and final provisions are embedded using the same sentence-transformer model (all-MiniLM-L6-v2), and cosine similarity matrices are computed between stakeholder embeddings and both draft and final embeddings. For units with an explicit legal reference from Step 2 (`explicit_reference = 1`), the script resolves the cited article/section/subsection against the draft provision hierarchy, handling edge cases such as missing section levels, subsection family matching, and cross-reference between notation formats. When only an article number is cited without finer location, the top 2 draft provisions within that article by cosine similarity are selected. Each explicit draft match is then propagated forward through the Step 1 continuity map to its corresponding final provision (`CONTINUITY_FROM_EXPLICIT_DRAFT`), but only when a forward mapping exists, since dropped draft provisions are excluded from continuity propagation. For all units (both explicit and non-explicit), the script adds the best semantic match to both draft and final provisions above a cosine threshold ($\tau = 0.35$ by default, raised to 0.65 for high-density policies). The best semantic final match is in turn projected backward through the continuity map to produce a `CONTINUITY_FROM_FINAL_MATCH` draft row, and the best semantic draft match is projected forward for a `CONTINUITY_FROM_DRAFT_MATCH` final row. Units whose best semantic match falls below the threshold receive a `NONE` row. The result is a comprehensive match table with six possible match sources per unit: `EXPLICIT_DRAFT`, `CONTINUITY_FROM_EXPLICIT_DRAFT`, `SEMANTIC_DRAFT`, `CONTINUITY_FROM_DRAFT_MATCH`, `SEMANTIC_FINAL`, and `CONTINUITY_FROM_FINAL_MATCH`. Five hard sanity checks and two soft warnings validate the output: schema/enum validation, explicit unit coverage, explicit draft set correctness (recomputed against the linking logic), final-row coverage for all units, and no continuity rows referencing dropped provisions.

Batching. Matched pairs are batched for GPT classification. The script derives a human-readable policy name from the regulation folder, filters out `NONE` matches and empty legal texts, deduplicates on (`stakeholder_id`, `version`, `stake_unit_id`, `legal_unit_id`, `match_source`), and builds token-aware batch files subject to both a hard count cap (`MAX_PAIRS_PER_BATCH = 12`) and a token ceiling

(MAX_PAYLOAD_TOKENS = 5,500, approximated at 4 characters per token). Each batch JSON file includes the policy name, task label, allowed classification labels, and the pair payloads.

Classification (Step 4). Each stakeholder–provision pair is classified using GPT-4.1 (temperature = 0) with OpenAI's Structured Outputs (strict JSON schema constraining output to pair_id + GPT_RELATION only). A proactive inter-batch delay of 20 seconds manages rate limits. Each pair is classified into one of five categories:

- **ALIGNED:** the stakeholder's position supports or is consistent with the legal provision
- **DIVERGENCE:** the stakeholder proposes an alternative formulation while engaging with the provision's substance
- **OPPOSITION:** the stakeholder demands deletion or fundamental rejection of the provision
- **NEUTRAL_OR_UNCLEAR:** the position is ambiguous or non-committal
- **NOT_SAME_ISSUE:** the stakeholder text and legal provision address different topics

DV scoring. The DV builder (dataset_builder.py) constructs draft–final correspondence by matching classified pair_ids across versions using the match_source chain, then applies the transition scoring matrix (Table A1).

Draft stance	ALIGNED	DIVERGENCE	OPPOSITION	NEUTRAL / NOT_SAME	Dropped
ALIGNED	0	−2	−2	−1	−2
DIVERGENCE	+2	0	0	0	0
OPPOSITION	+2	0	0	+2	+2
NEUTRAL_OR_UNCLEAR	+1	0	−1	0	0
NOT_SAME_ISSUE	0	0	0	0	0

Table A1. Transition scoring matrix mapping draft stance (rows) to final legislative outcome (columns).

Note: Scores range from +2 (maximum improvement: an opposed or divergent position becomes aligned in the final text) to −2 (maximum deterioration: an aligned position is dropped or reversed). NOT_SAME_ISSUE pairs score zero regardless of the final outcome, as they represent non-overlapping content.

For provisions newly introduced in the final text (no draft counterpart): ALIGNED = +1, OPPOSITION or DIVERGENCE = −1, other = 0.

Key rules: (1) Draft_dropped pairs are only scored when `match_source` $\in$ {EXPLICIT_DRAFT, CONTINUITY_FROM_FINAL_MATCH}, since semantic-only matches to dropped provisions score 0, since the stakeholder never specifically cited that provision and no outcome is attributable to them. (2) Article-level aggregation, in two steps. First, nesting is resolved: within an article, any two scored provisions standing in a containment relation (a section and a subsection within it, an article-level paragraph and any provision beneath it) are reduced to one, so that the larger magnitude is retained where the two share a sign, and both are discarded where their signs differ. Second, the surviving provisions, none of which contains any other, determine the article score on an extensity basis: the score is bounded to $\{-2, -1, 0, +1, +2\}$, an article whose surviving provisions point in opposite directions scores 0, and magnitude beyond ± 2 is not recorded (see Section 4.2 for why). The unbounded sum is retained in the output for reference but is not the dependent variable. (3) `dv_norm` = `total_relative_success` / `theoretical_max`, bounded $[-1, +1]$, where `theoretical_max` = $2 \times$ number of articles scored.

Aggregation: nesting resolution and the article score. Because positions are observed at several levels of the legal hierarchy, some scored provisions are nested inside others, since a stakeholder may be scored both on a section and on a subsection within it, and nested observations are not independent evidence of what the actor obtained, since the two may register one position at two magnitudes. Within each article, therefore, scored provisions standing in a containment relation are resolved before anything is aggregated. Where a nested pair carries the same sign, only the larger magnitude is retained, the smaller being treated as the same position observed less precisely. Where a nested pair carries opposite signs, neither is retained: a finer and a coarser reading of the same part of the text that contradict one another cannot both be evidence, and the pair contributes nothing. A provision scored zero carries no information and is dropped from the comparison rather than treated as conflicting. What survives is a set of provisions none of which contains any other.

Containment is determined positionally on the four-level unit identifier (version:article.section.subsection.subsubsection). Unit *a* contains unit *b* where the two share a version and article and, at every subordinate level, either match or *a* stands over a specific value in *b*, either because *a* covers the whole of that level or because *a*'s path terminates there. A level marked absent is not a wildcard spanning its siblings: a provision sitting directly under an article, outside any numbered section, is a sibling of the sections rather than their parent.

The surviving provisions then determine the article score, which is bounded to $\{-2, -1, 0, +1, +2\}$. An article whose surviving provisions point in opposite directions scores 0, and magnitude beyond ± 2 is not recorded; the unbounded sum is retained in the output for reference but is not the dependent variable. Article scores are summed to a total and normalised by $2 \times$ the number of articles engaged, giving $dv_norm \in [-1, +1]$. The rationale for bounding the measure is given in Section 4.2.

Precedent. The treatment of nested observations has no direct precedent in this literature, because no prior study of preference attainment observes positions at more than one level of textual granularity: where the unit of analysis is a consultation question or a survey item, units cannot contain one another and the question does not arise. It arises here as a consequence of the finer measurement rather than as a defect of it. A related move is nonetheless visible in the closest precedent, Bunea & Ibenskas (2017) recoding fine-grained positions into a smaller set of demands precisely "to account for the similarities among the positions."

A.4 IV Track: Dyadic Pair Generation, Classification, and Network Construction (Steps 5–6)

Pair generation (Step 5). Candidate stakeholder pairs are generated for dyadic comparison. Two stakeholder units are paired when they engage with the same legal provision, through either of two channels: (a) both units carry an explicit article citation (`explicit_reference` = 1 from Step 2) pointing to the same article, or (b) both are linked to the same `legal_unit_id` in the Step 3 match table (any version, any match source). Within-stakeholder pairs are excluded. A noise filter then retains only pairs where the two units share an explicit article citation or are linked to at least two shared legal units, discarding single-link semantic-only pairs that tend to produce uninformative comparisons. The upstream semantic similarity threshold from Step 3 ($\tau = 0.65$) ensures that all semantic matches reaching this step already exceed a minimum relevance floor (Table A2). For high-density policies, a per-dyad cap (`MAX_PAIRS_PER_DYAD` = 20) retains the most evidence-rich pairs, those sharing the most legal units, and discards the remainder. After deduplication (order-invariant: $A-B = B-A$), the surviving pairs are written into batch files of 50 pairs each. Each pair record carries a short opaque token (`pid_NNNNNN`) that GPT sees and echoes back, alongside the full canonical `pair_id` stored for downstream joining; this prevents GPT from truncating long compound identifiers.

Similarity band	Cosine range	N pairs	True continuity (%)	Assessment
Very high	0.90–1.00	~500	99%	Near-certain matches
High	0.80–0.90	~1,200	95%	Confident matches
Moderate	0.70–0.80	~800	82%	Threshold region
Marginal	0.60–0.70	~600	48%	Below threshold
Low	0.50–0.60	~400	15%	Rejected

Table A2. Empirical validation of the cosine similarity threshold ($\tau = 0.65$) used to filter semantic matches before dyadic pair generation.

Note: Each row reports a similarity band, the number of pairs, and the share classified as NEUTRAL, ALIGNED, or OPPOSED by the dyadic classifier across policies 1–9 ($n = 34,014$). Below 0.65, over 84% of pairs are NEUTRAL, contributing noise without informative signal. The 0.65 threshold balances recall of genuine substantive overlap against dilution by irrelevant pairings.

Classification. Each dyadic pair is classified using GPT-4.1 (temperature = 0) with OpenAI's Structured Outputs (strict JSON schema constraining output to pair_id + GPT_RELATION). The system prompt contains a playbook that instructs the model to apply a four-step decision procedure: (1) a same-issue gate that checks whether both units address the same concrete regulatory mechanism, returning NEUTRAL if they do not; (2) a normativity gate that checks whether both units contain a clear normative stance, returning NEUTRAL if either is purely descriptive; (3) stance typing that independently classifies each unit's position as support, oppose, or change; and (4) a label mapping that combines the two stance types into the final classification. Two units are ALIGNED when both support, both oppose, or both propose the same change to the same mechanism. They are OPPOSED when their positions are incompatible: support versus oppose, support versus change, or different changes to the same provision. They are NEUTRAL when neither the same-issue nor the normativity gate is passed, or when the relation cannot be confidently determined. The playbook explicitly warns against common failure modes: token-overlap traps (shared words do not imply shared mechanism), principle-versus-mechanism mismatches, and the assumption that "both want change" implies alignment. A proactive inter-batch delay of 20 seconds manages rate limits. Each batch result is validated against the input pair_ids (exact set match, no duplicates, no extras), and the short pid tokens are remapped back to canonical pair_ids before saving.

Network construction (Step 6). Pair-level classifications are aggregated into a signed dyadic edge list and stakeholder-level centrality is computed. NEUTRAL classifications are dropped (no edge). For each stakeholder dyad (i, j), the script

counts the number of ALIGNED and OPPOSED pair-level classifications, yielding $\text{tie_strength} = n_{\text{aligned}} + n_{\text{opposed}}$. The signed edge weight is $w_{\text{diff}} = n_{\text{aligned}} - n_{\text{opposed}}$, and the normalised weight is $w_{\text{norm}} = w_{\text{diff}} / \text{tie_strength}$. Stakeholder centrality is computed by summing incident dyadic weights: $C_{\text{signed_strength_norm}} = \sum_j w_{\text{norm}}(i, j)$ across all dyads involving stakeholder i , and $\text{degree} = \text{the count of unique dyads with tie_strength} > 0$. A label normalisation layer maps any variant spellings or legacy labels (DIVERGENCE, NOT_SAME_ISSUE, NEUTRAL_OR_UNCLEAR) to the canonical three-label set before edge computation.

A.5 Centrality Measures

Three centrality measures are computed over the signed stakeholder network of each policy.

Signed strength (positional convergence). For stakeholder i , $C_{\text{signed_strength_norm}} = \sum_j w_{\text{norm}}(i, j)$, summed over all other stakeholders j with whom i shares at least one engaged provision, where $w_{\text{norm}}(i, j) = (n_{\text{aligned}} - n_{\text{opposed}}) / (n_{\text{aligned}} + n_{\text{opposed}})$ across the provisions the pair both addressed. The measure is bounded by the number of an actor's counterparts and takes negative values where opposing ties predominate. It follows the signed-network centrality literature (Liu et al. 2020) in treating agreement and opposition as contributions of opposite sign to the same quantity, rather than as separate variables.

Degree. The count of an actor's ties irrespective of sign, following Freeman (1978/79). Degree measures how much of the field an actor engages, not how far it agrees with it, and is included because it operationalises the connectivity reading the framework rejects.

Betweenness. The normalised share of shortest paths between other pairs of nodes that pass through an actor (Freeman 1978/79), computed on the absolute-weight graph. Betweenness presupposes a flow model: a shortest path is consequential only if something traverses it (Borgatti 2005). It is included for the same reason as degree, to allow the flow reading to be tested rather than assumed, and its behaviour is discussed in Section 6.1.

A.6 Parameters and Thresholds

Legislative complexity is measured by five per-policy indicators: log provisions, the count of provisions in the draft; average depth, the mean hierarchical depth of provisions; the LIX readability index (Björnsson 1968); word entropy, the Shannon entropy of the unigram distribution; and the participation ratio, $PR = \text{Tr}(C)^2 / \text{Tr}(C^2)$ on the covariance matrix of provision embeddings, which is high where provisions span many semantic directions and low where they cluster in few (Roy & Vetterli 2007). Collinearity between the participation ratio, word entropy and log provisions precludes their simultaneous inclusion, so complexity interactions are estimated one indicator at a time.

[Table A3 here] — Consolidated pipeline parameters and thresholds. Reports the key hyperparameters governing each processing step, including cosine similarity thresholds, batch sizing limits, GPT model versions, and token ceilings. All parameters were set before the full-sample run and held constant across all 44 policies.

Parameter	Value	Step	Purpose
TAU_CONT_LEGAL	0.70	1	Draft–final continuity cosine threshold
TOPK_CAND	50	1	Blocked top-K candidates per provision
TAU_DRAFT_WEAK / TAU_FINAL_WEAK	0.35 (default)	3	Semantic match cosine threshold; raised to 0.65 for high-density policies
MAX_PAIRS_PER_BATCH	12	3	Stance classification batch size cap (hard count)
MAX_PAYLOAD_TOKENS	5,500	3	Stance classification token ceiling per batch
PAIRS_PER_BATCH	50	5	Dyadic classification batch size
MAX_PAIRS_PER_DYAD	20	5	IV pair cap (enabled for high-density policies)
APPROX_TOTAL_TOKEN_CEILING	12,000	2	Segmentation per-request ceiling (raised to 32,000 for long papers)
APPROX_TOTAL_TOKEN_CEILING	25,000	4	Stance classification per-request ceiling
APPROX_TOTAL_TOKEN_CEILING	20,000	5	Dyadic classification per-request ceiling
INTER_BATCH_DELAY_S	20	4, 5	Proactive rate-limit spacing (seconds)

Embedding model	all-MiniLM-L6-v2	1, 3	Sentence-level embeddings
GPT model	gpt-4.1-mini	2	Stakeholder segmentation and filtering
GPT model	gpt-4.1	4	Stance classification
GPT model	gpt-4.1	5	Dyadic relation classification
Temperature	0	2, 4, 5	Deterministic output

A.7 Verification and Known Limitations

The matching step (Step 3) runs seven automated checks after each policy: schema and enum validation, explicit-reference unit coverage, explicit draft set correctness (recomputed against the linking logic), final-row coverage for all units, and absence of continuity rows referencing dropped provisions. Cross-instrument citation noise from Step 2 (e.g., references to TFEU articles or the regulation's own preamble rather than the target regulation) is caught through post-processing audits that clear the article fields for units where the reference belongs to another legal instrument. Batch sizing in Step 3 uses conservative token estimates to avoid exceeding GPT context windows. All raw GPT outputs are archived as classified and annotated JSON files for replication and auditing, and the full prompts and playbooks used at each GPT step are available in the supplementary materials.

Three technical limitations should be noted. First, structural subdivisions below the 4-level hierarchy (e.g., sub-items within a subsection) are not parsed separately in Step 1, meaning that some fine-grained provisions are aggregated into their parent unit. Second, the GPT segmentation in Step 2 occasionally misassigns article references from cross-instrument citations despite explicit prompt instructions to exclude them. Third, the token estimation in Step 3 batching uses a 4-character-per-token approximation that slightly undercounts actual token usage.

A.8 Extended Regression Results

	(1) +Ctrl	(2) +Sect.	(3) C×Size	(4) C×PR	(5) Coal.	(6) Coal.+Q	(7) T1b	(8) T2a	(9) T3c
C_norm	+0.084** (0.032)	+0.085** (0.032)	+0.089** (0.034)	+0.086* (0.033)			+0.009*** (0.003)	+0.001* (0.001)	+0.008* (0.003)
Degree	−0.121*** (0.036)	−0.116** (0.036)	−0.105** (0.030)	−0.107** (0.030)			ctrl (0.030)	ctrl (0.030)	ctrl (0.030)
log(FTE)	−0.035 (0.030)	−0.042 (0.030)	ctrl (0.030)	ctrl (0.030)			ctrl (0.030)	ctrl (0.030)	ctrl (0.030)
log(Cost)	+0.019 (0.023)	+0.020 (0.023)	ctrl (0.023)	ctrl (0.023)			ctrl (0.023)	ctrl (0.023)	ctrl (0.023)
Brussels office	+0.012 (0.060)	+0.010 (0.060)	ctrl (0.060)	ctrl (0.060)			ctrl (0.060)	ctrl (0.060)	ctrl (0.060)
EP accredited	+0.064 (0.059)	+0.070 (0.059)	ctrl (0.059)	ctrl (0.059)			ctrl (0.059)	ctrl (0.059)	ctrl (0.059)
Non-EU HQ	−0.013 (0.066)	−0.017 (0.066)	ctrl (0.066)	ctrl (0.066)			ctrl (0.066)	ctrl (0.066)	ctrl (0.066)
Sector FE		Yes							
log(Provisions)			−0.046 (0.054)						
C_norm × Size			−0.055 (0.032)						
Partic. ratio				+0.088 (0.050)					
C_norm × PR				+0.054 (0.030)					
Coalition size					−0.060 (0.091)	−0.094 (0.094)			
Coalition cohesion					+0.120 (0.117)	+0.163 (0.119)			
Q (signed)						−0.112 (0.087)			
log(sub. chars)							+0.004 (0.006)		
chars × C_norm							−0.001** (0.000)		
is_umbrella								−0.007 (0.013)	
C_norm × openness									−0.002* (0.001)
N	1,966	1,966	1,709	1,709	212	212	1,709	833	1,966
ICC	0.090	0.039	0.092	0.087	0.099	0.100			
AIC	5,485	5,477	4,764	4,762	609	610			

Standard errors in parentheses. $p < 0.10$, * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$.
Models 1–2: full model progression (controls, sector FE). Models 3–4: complexity subsample (36 policies).
Models 5–6: coalition-level (212 coalitions). Models 7–9: robustness tests. "ctrl" = included but coefficients not shown.

Table A4. Extended regression results.

Note: Models 1–2: full model progression adding resource/access controls and sector fixed effects. Models 3–4: complexity interactions (regulation size × C_norm, participation ratio × C_norm) on the 36-policy subsample with available complexity data. Models 5–6: robustness tests for submission intensity and umbrella membership. "ctrl" indicates that the variable is included but coefficients are suppressed for space.

Appendix B. Language-Model Classification: Validity, Reliability, and Validation

Both the dependent and the independent variable are constructed by language-model classification. This appendix sets out the basis for treating those classifications as measurements, the design features that follow from it, the residual risks, and the validation protocol.

B.1 Why a language model rather than word-frequency scaling

Bunea & Ibenskas (2015) show that word-frequency scaling of the Wordfish family is largely invalid for European position documents. It forces a single dimension where the underlying space is multidimensional; it assigns spurious positions to texts that state none; it is contaminated by differences of author, register and purpose; and, most consequentially, it discards the technical vocabulary and the numbers, whether thresholds, dates or quantities, through which positions in this domain are actually stated. Classification that reads meaning in context avoids each of these failures, and does so at the level of the individual provision. The critique is specific to word-frequency methods, so the claim here is not that this approach overcomes their assumptions but that it is not subject to them; whether it succeeds is a matter for validation rather than assertion.

B.2 Evidence on validity

The strategy of using a language model to read texts for meaning and return a discrete position has been validated at scale. Benoit et al. (2026) estimate party positions on six issue dimensions from 235 manifestos in 21 languages and report correlations with expert-survey benchmarks of .87 to .92, at or near the ceiling that split-sample comparisons among the experts themselves establish. The relevant comparison is not with perfect accuracy but with the reproducibility of the human procedures such measures replace: trained Manifesto Project coders applying a fixed scheme achieve reliabilities between .3 and .5 (Mikhaylov et al. 2012), against inter- and intra-model agreement of roughly .90.

Two features of that design are shared here. Submissions are segmented into atomic units preserving the stakeholder's own wording before any stance is assigned, rather than passing whole documents to the model. This mitigates attention drift on long texts (Liu et al. 2023) and leaves a traceable intermediate representation that human review can inspect; Benoit et al. find that scoring directly from full documents returns substantially more missing values, so the step improves retrieval rather than

discarding information. And the classification admits categories for text that states no position or does not address the provision at issue. These are not residuals absorbing classification failure but a safeguard: a model not permitted to decline is a model forced to generate a position where the evidence for one is absent. That matters additionally here because the network is built only from provisions both stakeholders addressed, so "nothing in common" must be distinguishable from "in disagreement."

Two differences from that study should be stated plainly. Their reliability figures are properties of an ensemble of three models pooled across runs, whereas the classifications here come from a single model, so those coefficients are not transferable. And their design scores texts onto pre-specified bipolar scales, whereas the categories used here are not positions on a scale. The validation logic transfers; the reliability coefficients do not.

B.3 Reproducibility and model drift

Classification is run at temperature zero and raw model outputs are archived. Model drift remains a concern for any language-model measurement, since the underlying models are revised without notice. The available evidence bounds it: Benoit et al. (2026) re-ran their full analysis three months later on updated versions of the same models, recovering correlations of .92 to .98, and replicated it with three open-weight models, obtaining ensemble correlations of about .95 with the proprietary results. This qualifies rather than removes the concern, having been established for a different task on a different corpus.

B.4 Correlated measurement error

Using the same model for both variables raises a distinct risk. Classification errors do not bias regression estimates so long as they are not systematically correlated with the variables of interest (Heseltine & Clemm von Hohenberg 2024). That condition is not guaranteed when both the dependent and the independent variable are model-derived: if alignment were over-classified in particular text types, the association between positional convergence and success could be inflated. The two tasks classify structurally different constructs from different inputs, a five-category stakeholder–legal stance and a three-category stakeholder–stakeholder relation, which limits though does not eliminate the concern (Abdurahman et al. 2025; McLaren et al. 2026). The construction procedure follows recommended practice in codebook design, structured output and archived raw outputs (Halterman & Keith 2026; Ziems et al. 2024; Lin & Zhang 2025).

B.5 Human validation protocol

A stratified coding workbook of 194 stance-classification pairs and 100 dyadic-comparison pairs has been prepared for two independent PhD-level coders working blind. Results will be reported as a Cohen's κ matrix covering both human–human and human–model comparisons, with no third adjudicator.

Reporting the human–human figure alongside the human–model one is essential rather than incidental, because it establishes the difficulty ceiling of the task. Bunea et al. (2026), validating automated measures of comment quality against 1,800 human-coded snippets, report that the comparisons they automate are "hard for human analysts as well," so that disagreement between coders bounds what any classifier can be expected to reproduce. The same applies with more force here, where the categories are evaluative rather than descriptive. The sample is stratified by category, so κ should be read against the observed class prevalence rather than as a population estimate.

[Table B1 here] — Inter-rater reliability for language-model classifications. Cohen's κ for each coder pair (human–human, human–model) across the dependent variable (five-category stance classification, $N = 194$ stratified pairs) and the independent variable (three-category dyadic classification, $N = 100$ stratified pairs). Pending completion of the coding protocol.
