

The Source of the Majority Advantage in Legislative Effectiveness: Separating Cartel Power from Unified Government in U.S. State Legislatures

Kisoo Kim*

August 24, 2026

Abstract

Majority-party legislators are known to be more effective lawmakers than their minority-party colleagues; where that advantage is made is not. Cox-McCubbins cartel theory locates it in the majority's intrachamber agenda-control powers, and Anzia and Jackman derive its cross-sectional implication: the advantage should scale with how many of those powers a chamber's rules codify. An alternative locates it across the branches, in the co-partisan second chamber and governor that unified government supplies. Two regression discontinuity designs on a 1987-2024 panel of U.S. state legislators test both at the member level. Chamber-majority status raises a legislator's effectiveness score by roughly three-quarters of a standard deviation of that score across legislators, and the premium scales with the cartel-power index net of professionalism. A parallel design on unified state government finds that co-partisan control of the other chamber and the governorship adds no detectable premium beyond the member's own chamber majority. The majority advantage is made inside the chamber.

Keywords: state legislatures; legislative effectiveness; political parties; agenda control; regression discontinuity; separation of powers

*Email: kisoo@virginia.edu

1. Introduction

In 2009, the median Democratic member of Indiana's lower chamber introduced roughly fifteen bills and shepherded six through chamber passage; the median Republican introduced half as many and passed none. After the 2010 election, when Republicans took back the majority, the ranking flipped: the median Republican introduced twelve bills and shepherded four through passage, the median Democrat introduced seven and passed one. Much else changed with the chamber. Republicans already held the governorship and the Senate, so the election that gave them the House also gave them unified control of state government, and 2010 was a national Republican wave besides. A before-and-after contrast in one state therefore cannot say which of these made the median Republican more effective: majority status in her own chamber, the co-partisan Senate and governor through which her bills now passed, or the electoral tide that delivered both.

This pattern is not specific to a single legislature. In Congress, majority-party members consistently post higher individual-level legislative effectiveness scores than minority-party members (Volden and Wiseman, 2014), and the same gap appears on the state legislative effectiveness scores of Bucchianeri et al. (2025), though its size varies markedly across chambers. In our data that gap follows a chamber's current majority, not a party or its legislators. Members of a party that has just won control post the higher scores in the very next session, and members of the party that has just lost it no longer do. Two institutional accounts locate the source of that advantage in different places. Cartel theory locates it inside the chamber: the majority caucus, working through co-partisan committee chairs and a co-partisan rules committee, decides which bills receive hearings and reach the floor and steers referrals and favorable reports toward its own members (Cox and McCubbins, 1993, 2005). Because every such tool operates bill by bill, a chamber whose standing rules codify more of them should confer a larger premium on each majority member, the member-level form of the comparative static Anzia and Jackman (2013) propose and test across state chambers. A cross-branch account locates part of the

advantage beyond the chamber: under unified state government, one party holding the governorship and both chambers (Kirkland and Phillips, 2020), a majority member's bills face a co-partisan second chamber and a co-partisan governor, so holding those institutions should add a premium of its own on the outcomes they decide. We ask where the majority advantage is made, and whether the part made inside the chamber scales with the cartel power the chamber's rules confer.

Prior work has approached the two loci one channel at a time, in part because close-election designs on partisan control identify a compound treatment bundling majority status with partisanship (Alpino and Crispino, 2024), and in state settings unified government as well. The within-chamber channel has been identified causally: McGrath and Ryan (2019) use a chamber-majority regression discontinuity in state legislatures to show that majority status delivers a committee-seat bonus and ideological stacking, and Napolio and Grose (2022) trace majority-control shifts onto legislator agenda behavior in the U.S. Senate, but on committee and roll-call outcomes rather than a member-level effectiveness measure, and without a cross-branch channel. The cross-branch channel likewise: Repetto and Sosa Andrés (2023) use a divided-government regression discontinuity to recover effects on individual legislators' ideology and on gubernatorial vetoes, but not effectiveness, and with chamber flips bundled into the unified-government transition rather than separately identified. Closest to our cartel theory test, Howard and Provins (2026) connect Cox-McCubbins and Anzia-Jackman procedural rules to state legislative effectiveness, but in cross-sectional ordinary least squares on party-aggregate, stage-specific outcomes, with unified government entering only as a control. No published study, to our knowledge, identifies both channels at the individual-member level on a common legislative effectiveness outcome and tests their separability.

The U.S. states make that test possible and dictate its form. The ninety-one partisan chambers we analyze, spread across forty-nine states, supply both the cross-sectional variation in cartel rules and the many close chamber-control elections that an individual-

member discontinuity design requires. But in the states the two loci move together: a close-election design on unified government carries the member's own chamber majority across its cutoff whenever her chamber is the institution closest to flipping, so the two questions have to be answered together. We therefore run two regression discontinuity designs in parallel on a 1987–2024 panel of state legislators. The first is a chamber RDD on the smallest uniform statewide vote shift that would flip chamber control: because the shift aggregates many district races, no actor can place a chamber precisely at the threshold, so members just inside it are as-if-randomly assigned to majority status, and a tide like Indiana's 2010 wave is balanced across the cutoff rather than confounded with it (Lee, 2008). The second is a state RDD on the smallest uniform shift that would flip unified-government status, with one modification to the Kirkland-Phillips design: where they treat unified government as any party holding all three institutions, appropriate for their state-policy outcomes, we re-orient the treatment to whether the member's own party holds the full coalition, because the any-party formulation pools a member whose party holds unified government with a member locked out by it and cancels the contrast. Both forcing variables are constructed deterministically from observed vote shares, without Monte Carlo simulation. Because the state cutoff still carries the member's own chamber majority for part of the near-threshold sample, we decompose its estimate by the institution that binds, and we read it against the two implications that separate a cross-branch premium from the chamber premium: which outcomes it should amplify, and that it should not respond to how richly one chamber's rulebook codifies cartel power.

The chamber-majority premium on the legislative effectiveness score is 0.82 points, roughly three-quarters of the score's standard deviation across legislators, and it scales across the cartel-power index as the theory predicts: the per-member premium rises by about 0.16 points with each additional codified tool, a gradient that survives a professionalism horse race and small-cluster-robust inference. The cross-branch design tells a different story from the one its primary contrast suggests. Members whose party holds

unified state government post a higher score than members whose party does not, by about a quarter of the chamber premium, but the cutoff is a compound margin: for one-fifth of the near-threshold sample the binding institution is the member's own chamber, so crossing it flips chamber majority together with unified status. Decomposed by binding institution, the contrast is the chamber premium surfacing in those cells. Where the governorship is the binding margin, so that neither chamber can flip, unified control of the remaining branches adds nothing distinguishable from zero on the effectiveness score or on any funnel outcome, with an upper bound of about an eighth of the chamber premium; the same design's gradient on the cartel-power index and its amplification profile across outcomes are the chamber pillar's, scaled by the own-chamber share. The intrinsic between-party effectiveness gap, conditional on both channels, is comparatively small.

Two contributions follow. First, the scaling hypothesis Anzia and Jackman propose at the chamber level passes its member-level test. Cartel theory has been tested mainly on variation over time across U.S. House sessions and speakers (Cox and McCubbins, 2005) and, cross-sectionally, on aggregate roll rates across state chambers (Anzia and Jackman, 2013); we provide the member-level, discontinuity-identified test on individual productivity. The gradient is large, robust under cluster-robust and wild-bootstrap inference, and survives the professionalism horse race. The contribution is not that cartel power appears in the states, which the aggregate roll-rate tests established, but that the scaling prediction holds at the level of individual-member productivity, identified off close elections.

Second, the majority advantage is made inside the chamber. The closest federal evidence, Volden and Wiseman's finding that the majority-party effectiveness advantage is about the same under unified and divided government, is descriptive; the state design identifies the same null off close margins and bounds it at an eighth of the chamber premium. Reaching that result requires reading the Kirkland-Phillips state-government design at the member level, which takes two steps the original does not need: treatment

must be party-specific, because the any-party definition pools the member whose party holds unified government with the member locked out by it and cancels the contrast to zero; and the estimate must be decomposed by binding institution, because the party-specific cutoff carries the member's own chamber majority for one-fifth of the near-threshold sample. Researchers extending state-government discontinuity designs to member-level outcomes (roll-call behavior, distributive returns, electoral safety) should expect both steps to matter: the first moves the contrast from zero to a quarter of the chamber premium, the second moves it back.

2. Theory and Institutional Channels

2.1 The chamber cartel channel

Cox and McCubbins (1993, 2005) develop the canonical account of why majority-party members enjoy a productivity advantage inside a legislative chamber. The mechanism is agenda control, exercised asymmetrically. The majority caucus, through its allied committee chairs and rules committee, withholds hearings, markups, and floor time from bills that would split the caucus or advantage the minority. Because the gatekeepers are majority members and the bills screened out disproportionately the minority's, the apparatus leaves majority members the larger share of legislative opportunities (Cox and McCubbins, 2005) — an allocation *Legislative Leviathan* traces to majority control over jurisdictions, chair scheduling power, hearing and markup rights, and staff (Cox and McCubbins, 1993). Legislative entrepreneurship is therefore likelier to yield a hearing, a markup, and a floor vote when the entrepreneur sits inside the caucus than outside it. The theory is institutional: it ties the size of the majority advantage to the rules of the chamber rather than to the ideologies or skills of the legislators in it — in contrast to conditional-party-government accounts, where the majority's agenda leverage is granted and used only when member preferences are homogeneous enough to sustain it (Ro-

hde, 1991; Aldrich, 1995) — and it makes predictions about how chamber rules shape the premium. Quasi-experimental evidence that majority agenda power operates in the state-legislative setting comes from Cox et al. (2010), who show that stripping the majority’s agenda control — Colorado’s 1988 GAVEL initiative — or switching it off across bill pathways within a chamber — California’s Suspense File — reshapes roll-call coalitions and the direction of policy movement; their outcomes are bill fates and roll-call patterns rather than individual-member productivity.

We expect that chamber-majority status produces a substantively large premium on individual-member effectiveness, operating through these intrachamber cartel mechanisms, and that the premium concentrates on outcomes decided inside the chamber — most directly, whether a member’s bills pass the home chamber. The scaling hypothesis requires one more step. Cox and McCubbins developed and tested the theory on the U.S. House, where the cartel’s toolkit is essentially a constant of postwar chamber organization; the comparative static we test — a larger per-member premium where the chamber’s standing rules codify more cartel power — is the member-level counterpart of the chamber-level comparative static Anzia and Jackman (2013) propose and test for the cross-section of state chambers. The member-level version follows from the incidence of the tools themselves: every codified blocking tool operates bill by bill, so each tool the rules add enlarges both the set of minority-member bills the majority can stop and the set of majority-member bills it can advance, widening the gap between an individual majority member’s realized productivity and her minority counterpart’s. Anzia and Jackman’s coding of chamber rules provides the operationalization. They code four chamber-rule binaries — committee denial of a hearing, committee denial of a report, a majority leader setting the floor calendar, and a majority-appointed committee setting it — together with the two summary indicators they build from these, any committee gatekeeping and any majority calendar control. We sum all six into a cartel-power index; because two summands are disjunctions of the other four, the index orders chambers by institutional in-

tensity rather than counting independent tools. At the top of that ordering standing rules concentrate agenda control sharply in the majority caucus; at the bottom it diffuses across actors who may or may not align with the majority on any given bill. The scaling hypothesis is that the size of the majority premium should track this institutional intensity.

A second testable implication concerns the boundary of the channel. Because the relevant agenda authority is intrachamber, the channel should act primarily on chamber-floor outcomes — whether a member’s bills pass the home chamber — and only indirectly on terminal-law outcomes that require the other chamber and the governor: a chamber cartel can move its majority’s bills through passage, but it cannot guarantee that they will agree. Section 4.2 tests the resulting channel-separation prediction.

2.2 The cross-branch coordination channel

A second institutional mechanism operates across the branches rather than within one chamber. Whether unified government raises aggregate lawmaking output is contested — Coleman (1999) finds unified governments more productive on significant legislation and more responsive to public mood, against a gridlock literature finding little unified-divided difference in major enactments — and the debate has stayed at the aggregate level. At the individual-member level, the closest federal benchmark finds the majority-party effectiveness advantage essentially unchanged under unified versus divided government — roughly one effectiveness point under both (Volden and Wiseman, 2014). The comparison is descriptive; we know of no federal study that design-identifies a member-level cross-branch premium. Kirkland and Phillips (2020) build a regression discontinuity design around close-to-unified state governments — configurations in which the same party holds the governorship and both chambers. Their substantive interest is in state-level policy, and they treat the design’s running variable as a measure of how close a state-year is to flipping between unified and divided government regardless of which party would hold the unified configuration. When a single party simultaneously controls

the governorship, the lower chamber, and the upper chamber, the path from a member's bill introduction to a signed law passes through allied institutions at every step: a committee chaired by a co-partisan, a floor controlled by a co-partisan caucus, an other-chamber vote in a co-partisan caucus, and a signature from a co-partisan governor.

We expect that cross-branch unified government produces a distinct premium on member effectiveness, one that operates through the full lawmaking pipeline rather than through intrachamber agenda control alone. The cross-branch channel should therefore concentrate on terminal-law outcomes: bills that become law, including the subset that are substantively significant. It should also be smaller than the chamber cartel premium: the chamber cartel premium is the effect of holding one institution, the cross-branch premium is the marginal effect of the remaining institutional alignments conditional on holding the chamber, and the marginal contribution of each subsequent alignment is plausibly less than the contribution of the first.

The substantive distinction between any-party and party-specific unified government matters here. From the perspective of a single member, the relevant question is whether her own party holds the three institutions, not whether some party does. A state-year in which the opposite party holds all three is, for a minority-party member, the worst possible institutional configuration — every veto point held by the opposing team — yet the any-party classification pools it with state-years in which the member's own party holds unified government, the best possible configuration. Our second pillar therefore treats unified government as party-specific: the running variable measures the smallest uniform vote shift that would flip the member's-party-unified status, and the treatment indicator is one only when the member's own party holds all three institutions.

2.3 Channel-separation predictions

The two accounts can be tested against the same data through two implications that follow from where each channel acts rather than from the data. The first concerns the menu

of outcomes. Because the chamber cartel channel acts on intrachamber decisions and a cross-branch channel on the full pipeline, a genuine cross-branch premium should be over-amplified on terminal-law outcomes relative to its own pillar’s effectiveness baseline, and the chamber premium on chamber-floor outcomes. A state design whose contrast were only the chamber premium carried across its cutoff would instead reproduce the chamber pillar’s amplification profile outcome for outcome.

The second concerns heterogeneity. The Anzia-Jackman coding operationalizes the chamber’s internal cartel features, so the chamber cartel premium should track the index. A cross-branch coordination premium should not, because what matters for coordination is whether the member’s party holds all three institutions, not how richly one chamber’s rulebook codifies cartel power; and a state-design contrast that were only the chamber premium would show the chamber gradient scaled down by the share of near-threshold cells in which the member’s own chamber is the binding institution. Sections 4 and 5 read the state pillar against both benchmarks, a distinct cross-branch channel and no cross-branch channel at all.

3. Data and Research Design

Our analytical panel covers state legislators in forty-nine states — ninety-one of the ninety-eight partisan chambers — over the 1987-2024 biennia: Nebraska’s nonpartisan unicameral falls outside the design, and seven chambers elected under multimember-district structures do not yield the district-level vote margins the forcing variables require. The full panel contains 94,107 member-terms; the close-margin analytical samples we use for identification are subsets of this within the primary bandwidth. The outcome measure is the State Legislative Effectiveness Score (SLES, hereafter LES) developed by Bucchianeri et al. (2025), in its 2026 per-state release, on the methodology of Volden and Wiseman (2014), whose effectiveness-score framework has a formal-theoretic

foundation in spatial models of lawmaking (Hitt et al., 2017). LES aggregates each member-term’s contribution across a five-stage funnel — bills introduced, bills receiving action in committee, bills receiving action beyond committee, bills passing the home chamber, bills becoming law — with weights that emphasize the rarer later stages and that distinguish substantively significant from minor legislation. We report results on LES itself and on its four most informative funnel components: `bill_pass` (bills passing the home chamber), `bill_law` (bills becoming law), `ss_pass` (substantively significant bills passing the home chamber), and `ss_law` (substantively significant bills becoming law). The funnel structure of LES is what allows the outcome-stage prediction of section 2.3 to be tested: chamber-floor passage and terminal lawmaking are different points on the same legislative pipeline, and the substantively significant subsets are higher-stakes counterparts.

The forcing variables for both pillars are constructed deterministically from observed vote shares: district-level legislative returns from the Klarner State Legislative Election Returns database, which extends back to 1967, joined for the state RDD with the gubernatorial margins described below. For Pillar A, the chamber RDD, the running variable is the smallest uniform statewide vote-share shift that would flip chamber majority control. For Pillar B, the state RDD, the running variable is the smallest uniform statewide shift that would flip the member’s party-specific unified-government status, computed over the three institution margins (governorship, lower chamber, upper chamber). The two sides of the cutoff are not symmetric in construction, because losing unified government requires losing only one institution while attaining it requires winning them all: for a member whose party currently holds unified government, the flip distance is the minimum across the three institution-specific margins, while for a member whose party does not, it is the maximum across the margins of the institutions her party does not hold. One consequence of the three-institution margin matters for interpretation, and section 4.1 addresses it directly: in the near-threshold sample the binding institution is the member’s

own chamber for roughly one-fifth of observations, so for that subset crossing the Pillar B cutoff flips chamber-majority status together with unified status. This approach follows Lee (2008) in treating close electoral margins as the identifying source; Pillar B's closed-form multi-margin construction parallels the multi-dimensional regression discontinuity design of Chin and Shi (2025).

The governor data require splicing. The Kirkland-Phillips replication archive provides clean gubernatorial vote shares through 2012; for 2013-2022, we splice in Dave Leip's Atlas of U.S. Presidential Elections gubernatorial data. On the 2008-2012 overlap window the two sources agree on every winner; their two-party shares differ by a standard deviation of 0.2 percentage points (median zero, largest 1.3). Pillar B's analytical sample excludes seven states (Arizona, Idaho, North Dakota, New Jersey, South Dakota, Washington, West Virginia) that have zero coverage on the construction because their lower or upper chamber uses multimember-district or jungle-primary structures that the Klarner database does not parse into smallest-shift-to-flip terms. This is the same exclusion that propagates through Pillar A's chamber RDD subsample, and a placebo test in the appendix confirms that the seven excluded states behave statistically like a random seven-state loss from the full set.

Both pillars use a signed running variable that crosses zero at the threshold: positive values indicate the treated state (the member's party holds chamber majority for Pillar A; the member's party holds unified state government for Pillar B), negative values indicate the control state, and zero indicates the boundary case, which we exclude. The primary bandwidth is five percent of statewide vote share, yielding analytical samples of 27,433 member-terms for Pillar A and 32,940 for Pillar B. The narrow-bandwidth choice trades sample size for identification cleanliness: at five percent the as-if-random assumption is most defensible, and the wider-bandwidth alternatives in the appendix admit observations farther from the close-margin condition. Both pillars use additive state-chamber-by-year (Pillar A) or state-by-year (Pillar B) fixed effects, with three member-level controls

(seniority, gender, ideology), a party indicator whose coefficient we report as the residual partisan gap, and no interaction terms in the primary specification. The estimator is local-linear inside the fixed bandwidth: a separate linear term in the running variable on each side of the cutoff, a triangular kernel that down-weights observations away from it, and the treatment coefficient read as the intercept shift at the threshold. A within-window difference in means would carry first-order bias from the slope of the outcome in the running variable, and here that slope is built in: both running variables are chamber- or state-level states with high period-to-period persistence, so the share of members whose party also held the status in the previous term rises with distance from the cutoff on both sides; a linear term on each side removes what a mean difference absorbs. The data-driven bandwidth and robust bias-corrected inference of Calonico et al. (2014) enter as a robustness check rather than the primary, because the mass points of a chamber-level running variable make its clustered variance estimator singular (section 6.5).

Standard errors require attention. With moderate cluster counts — sixty-seven state-chambers for Pillar A and forty-one states for Pillar B in the analytical samples (forty-two states in the population panel) — the asymptotic CR1 may under-cover. We therefore report three inference methods side by side on every main result: the primary cluster (state-chamber for Pillar A, state for Pillar B), a secondary cluster (state, state-year), and a wild cluster bootstrap on the primary cluster with 999 Rademacher replicates and restricted residuals (Cameron et al., 2008). Conclusions in the main text rely on agreement across all three.

The party-specific re-orientation of Pillar B is the first of two steps that reading the state-government design at the member level requires; the second, the decomposition of its compound cutoff by binding institution, is in section 4.1. Kirkland and Phillips (2020) treat unified government as an any-party condition, appropriate for their state-policy outcomes: from the perspective of policy outputs, what matters is whether some unified coalition exists to enact policy. Our individual-member outcome requires a party-

specific construction in which a member observation is treated only when her own party holds all three institutions. The opposite case — when the opposing party holds unified government — is control. Pooling these two cases under one “treated” label would mix substantively opposite institutional conditions; section 6.3 tests the re-orientation directly: pooling the two cases drives the contrast to zero. The two forcing variables are only moderately correlated (approximately +0.4 on the overlap window with the Kirkland-Phillips construction), consistent with the substantive distinction: the party-specific running variable carries party-identifying information that the any-party construction collapses.

A short validity battery closes this section, with the appendix carrying the full detail. Covariate continuity tests on six predetermined characteristics (Squire professionalism, seniority, gender, ideology, chamber size, appointment status) pass five of six on the chamber pillar and three of six on the cross-branch pillar. The failing cells are appointment status on the chamber pillar (0.10 standard deviations, $p = 0.02$), gender and chamber size on the cross-branch pillar (0.09 and 0.01 standard deviations, significant only because of the sample size), and Squire professionalism on the cross-branch pillar, which fails on magnitude (0.16 standard deviations, $p = 0.20$); the state fixed effects absorb part of that imbalance, and section 5 shows professionalism does not moderate the cross-branch contrast. A density-continuity check on the running variable, in the spirit of McCrary (2008), compares observation counts on either side of the threshold at three windows; it passes at all three for the chamber pillar and at two of three for the cross-branch pillar, where the failing intermediate ($\pm 1\%$) window is bounded by passing tests at $\pm 0.5\%$ and $\pm 2\%$. A joint subset that imposes both pillars’ close-margin conditions (15,542 member-terms) reproduces the chamber premium within seven percent and raises the cross-branch contrast by two-fifths, the direction section 4.1’s decomposition predicts (Appendix G).

A pre-LES placebo, which regresses the previous term’s effectiveness score on current treatment, returns 0.28 on the chamber pillar (standard error 0.10) and 0.12 on the cross-

branch pillar (standard error 0.08, not distinguishable from zero). Both coefficients are the size that status persistence predicts rather than a sign of anticipation: majority and unified status persist from one term to the next, so members just inside the cutoff today were disproportionately inside it last term, and last term's premium appears in last term's score. The direct test is a covariate discontinuity in lagged status: at the chamber cutoff the probability that the member's party held the majority in the previous term jumps by 0.34 (standard error 0.13), and the chamber premium times that jump, 0.27, reproduces the placebo to within 0.01; at the state cutoff the lagged-unified jump is 0.18 (standard error 0.11), and the placebo is within its own noise. Chamber-terms where the majority actually flipped return the mirror image of the premium, -0.87 , as a genuine premium must and anticipation would not. Appendix I reports the diagnostics for both pillars and bounds the effect of the imbalance on the primaries at about 0.01.

4. Results — Two Channels

4.1 Primary results

Consistent with the chamber cartel channel of section 2.1, chamber-majority status is worth roughly three-quarters of a standard deviation of legislator effectiveness: the gap between otherwise comparable members just inside and just outside a chamber majority is 0.82 LES points (standard error 0.11), against a standard deviation of 1.10 in the analytical sample (1.13 in the full panel). Table 1 reports the estimate across five outcomes and three standard-error methods. The four funnel components scale with their base rates: the largest absolute effect is on `bill_pass`, the chamber-floor outcome that Cox and McCubbins identify as the most direct channel of cartel agenda control. The estimate is robust across the multi-SE machinery — the primary state-chamber cluster, the secondary state cluster, and the wild cluster bootstrap all yield p-values below 0.001 — and the wild bootstrap symmetric percentile-t confidence interval lies entirely above zero (Table 1).

Table 1: Primary RDD estimates of the two institutional channels

	Pillar A — chamber cartel ($\hat{\mu}$)	Pillar B — cross-branch ($\hat{\gamma}$)
<i>Panel A. LES, three inference methods</i>		
Primary cluster SE	0.818 (0.108)***	0.198 (0.059)***
Secondary cluster SE	0.818 (0.132)***	0.198 (0.045)***
Wild bootstrap 95% CI	[0.458, 1.192]	[0.054, 0.333]
<i>Panel B. Five outcomes (primary cluster)</i>		
LES	0.818 (0.108)***	0.198 (0.059)***
bill_pass	4.509 (1.040)***	0.698 (0.409)
bill_law	2.785 (0.585)***	1.154 (0.309)***
ss_pass	0.270 (0.083)**	0.051 (0.143)
ss_law	0.169 (0.059)**	0.129 (0.127)

Note: Member-term local-linear RDD estimates (separate slopes on each side of the cutoff, triangular kernel) within the $\pm 5\%$ primary bandwidth (signed running variable, tied cases excluded), 1987–2024; standard errors in parentheses. Pillar A primary cluster is state-chamber ($N = 27,433$; 67 clusters) and the secondary cluster is state; Pillar B primary cluster is state ($N = 32,940$; 41 clusters) and the secondary cluster is state-year. The wild cluster bootstrap reports the symmetric percentile-t 95% interval (Cameron–Gelbach–Miller, 999 replications). Stars: *** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$. The intrinsic Democratic–Republican gap recovered alongside each channel is comparatively small ($\hat{\delta} \approx -0.15$ to -0.21 , standard errors 0.06; see section 4.1 for its identification status).

Members whose own party simultaneously holds the governorship and both chambers post higher effectiveness than members whose party does not: the primary estimate of the unified-government contrast is 0.20 LES points (standard error 0.06), distinguishable from zero under all three inference methods (Table 1). Read on its own, the contrast would suggest the cross-branch coordination channel of section 2.2 operating at about a quarter of the chamber premium. It cannot be read on its own.

The Pillar B cutoff is a compound margin, and the primary estimate must be read as such. Because the running variable measures the distance to flipping unified status across three institutions (section 3), what flips at the cutoff depends on which institution is binding, and Table 2 decomposes the estimate accordingly. In the near-threshold sample the governorship is the binding margin for 58 percent of member-terms, the member’s own chamber for 20 percent, and the other chamber for 22 percent, and the probability of chamber-majority status jumps by 0.18 at the cutoff (Table 2, row 2). For the own-

chamber-binding subset, crossing the cutoff flips chamber-majority status together with unified status, and that subset's contrast is 0.67, indistinguishable from the chamber cartel premium. For the governor-binding subset, where neither chamber can flip at the margin, the contrast is -0.08 (standard error 0.06; wild bootstrap interval -0.22 to 0.07): this is the design's clean cross-branch estimate, and it is not distinguishable from zero. The other-chamber-binding subset gives -0.05 . Three further quantities agree that the primary contrast is the chamber premium surfacing through the own-chamber-binding cells rather than a coordination premium that happens to be small. The majority-status jump times the own-chamber contrast, 0.18 times 0.67 , accounts for 0.12 of the 0.20 ; conditioning the contrast on majority status directly leaves 0.08 (Table 2, last row), a regression adjustment for a covariate that itself jumps at the cutoff and therefore a bound rather than a discontinuity estimate. Section 4.2 shows that the cross-branch estimates scale across the funnel outcomes in the same proportions as the chamber estimates, and section 5.2 that the cross-branch pillar's gradient on the cartel-power index equals the chamber gradient times the own-chamber share. We therefore report the cross-branch coordination premium as not distinguishable from zero on the effectiveness score, with an upper bound of about 0.1 LES points, an eighth of the chamber premium, and read the primary 0.20 as the contrast for the bundled configuration the compound cutoff delivers: a member whose unified government hangs on her own chamber's majority is receiving the cartel premium, not a coordination premium on top of it.

Unified government is also historically less common than chamber majority. Across the 1987-2024 panel, roughly one quarter of valid state-year-party cells correspond to the member's-party-unified configuration, with divided government in some form the modal arrangement. The cross-branch contrasts are identified off the close-to-flip subset within this rarer configuration, and the estimates in Table 2 are average causal effects at that boundary rather than state-population averages.

The intrinsic between-party effectiveness gap — the residual Democratic-Republican

Table 2: Inside the cross-branch estimate — the compound cutoff decomposed

Quantity	Estimate (SE)	Wild 95% CI	N	G
Primary unified-government contrast ($\hat{\gamma}$, Table 1)	0.198 (0.059)	[0.054, 0.333]	32,940	41
Discontinuity in majority status at the Pillar B cutoff	0.177 (0.071)	[−0.023, 0.357]	32,940	41
Governor-binding subsample (clean cross-branch)	−0.075 (0.060)	[−0.218, 0.066]	18,939	40
Own-chamber-binding subsample	0.672 (0.117)	[0.357, 1.023]	6,650	31
Other-chamber-binding subsample	−0.051 (0.111)	[−0.292, 0.210]	7,351	31
Unified-government contrast, conditioning on majority status	0.083 (0.029)	[0.024, 0.137]	32,940	41

Note: All rows use the Pillar B primary specification (state and year fixed effects, state cluster, $\pm 5\%$ bandwidth, tied cases excluded); wild cluster bootstrap as in Table 1. The binding institution is the margin attaining the minimum (unified side) or the maximum over unheld institutions (non-unified side) in the section 3 construction; near-threshold composition: 57.5% governor, 20.2% own chamber, 22.3% other chamber. Row 2 replaces the outcome with the member’s majority-status indicator. The final row controls for majority status directly; because unified status entails own-chamber majority for over 99 percent of unified member-terms, the two are nearly collinear and the conditioned coefficient is a bound, not a clean decomposition — the binding-institution split is the credible route.

difference after both institutional channels are accounted for — is comparatively small: on the effectiveness score, roughly -0.15 in the chamber pillar and -0.21 in the state pillar (standard errors 0.06). Two caveats bound this quantity: unlike the two institutional premia it is not discontinuity-identified — it is the coefficient on the member’s party indicator, a covariate-adjusted association we read as descriptive — and it is statistically

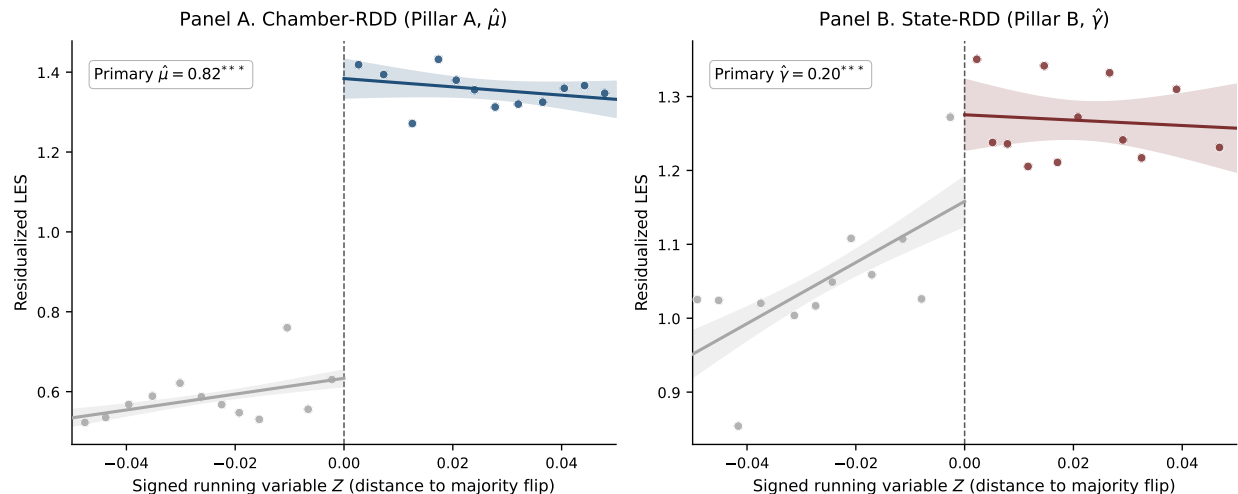


Figure 1: Two-pillar RDD binscatter. Binned means of covariate-residualized LES against the signed running variable within the $\pm 5\%$ bandwidth, with a local-linear fit and 95% confidence band on each side of the threshold. Panel A is the chamber RDD (Pillar A); Panel B is the party-specific state RDD (Pillar B). The boxed quantity in each panel is the primary local-linear RDD estimate (Table 1).

distinguishable from zero, not a null. The substantive reading is therefore a claim about relative magnitudes: the apparent partisan effectiveness gap in U.S. state legislatures is produced mainly by who holds the institutional levers, not by intrinsic differences in how Democratic and Republican legislators approach the work of lawmaking. This is consistent with the nationalization of state politics that Hopkins (2018) documents — national-party identity has become increasingly load-bearing in state elections — adding an institutional layer his account does not address: how that identity translates into legislative output runs through state-level institutional configurations.

Figure 1 shows the binscatter RDD plots for both pillars: binned means of covariate-residualized LES against the signed running variable, with a local-linear fit and shaded 95% confidence band on each side of the threshold, Panel A for the chamber RDD and Panel B for the party-specific state RDD, the primary estimate (Table 1) boxed. The residualized level steps up sharply at the threshold in Panel A; in Panel B the step is small and the bands overlap, which is what a bundled contrast of 0.20 whose clean component is near zero looks like.

4.2 Channel separation across outcomes

Section 2.3 predicts which outcomes each channel should amplify. The chamber cartel channel acts on intrachamber decisions, so chamber-floor outcomes (`bill_pass`) should be over-amplified under chamber RDD identification relative to that pillar's own LES baseline. A cross-branch coordination channel acts on the full lawmaking pipeline, so terminal-law outcomes (`bill_law`, `ss_law`) should be over-amplified under state RDD identification relative to that pillar's own LES baseline. Table 3 reports the head-to-head comparison: for each outcome, the chamber cartel estimate, the cross-branch estimate, each estimate's ratio to its own pillar's LES estimate, and the difference of those ratios. A positive difference favors the cross-branch channel; a negative difference favors the chamber cartel channel. Because each ratio is taken within a pillar, the cross-pillar comparison is of unitless amplification profiles. If the cross-branch pillar were carrying only the chamber premium, the two profiles would coincide and every difference would be zero, so the table also reports a pairs cluster bootstrap interval for each difference and a joint test that the four differences are zero.

The four differences carry the signs the prediction assigns, but none is distinguishable from zero, and the joint test does not reject equal profiles ($p = 0.83$). The one sizable difference is on `bill_law`, where the cross-branch ratio is 5.8 against the chamber pillar's 3.4, but its interval runs from -0.2 to 10.4 because the cross-branch LES estimate in every denominator is itself 0.20 with a standard error of 0.06 . The evidence is consistent with equal amplification profiles and, given the width of the intervals, with a wide range of unequal ones; it does not show that the cross-branch pillar amplifies terminal-law outcomes. Decomposing each outcome by binding institution (Appendix P) sharpens the reading: in the governor-binding cells the cross-branch estimate is at or below zero on all five outcomes, `bill_law` included (-0.07 , standard error 0.42), while the own-chamber-binding cells carry the pooled `bill_law` (3.42) and `bill_pass` (3.86) estimates, the chamber funnel. Under a local-constant estimator the same governor-binding `bill_law` cell is positive (0.71 ,

Table 3: Channel separation across outcomes (head-to-head)

Outcome	$\hat{\mu}$ (Pillar A)	$\hat{\gamma}$ (Pillar B)	μ ratio	γ ratio	Ratio diff ($\gamma - \mu$)	95% interval
LES	0.818	0.198	1.00	1.00	—	(baseline)
bill_pass	4.509	0.698	5.51	3.52	−2.00	[−7.69, 2.62]
bill_law	2.785	1.154	3.41	5.81	+2.41	[−0.17, 10.43]
ss_pass	0.270	0.051	0.33	0.26	−0.07	[−1.62, 1.53]
ss_law	0.169	0.129	0.21	0.65	+0.45	[−0.58, 2.26]

Note: Each pillar’s ratio is the outcome estimate divided by that pillar’s own LES estimate; the ratio difference is the cross-branch ratio minus the chamber cartel ratio, positive where the cross-branch channel over-amplifies the outcome relative to its baseline. The interval is the 95% percentile interval from a pairs cluster bootstrap over states (999 replications) re-estimating both pillars on each draw; the joint Wald test that all four differences are zero gives $\chi^2(4) = 1.46$, $p = 0.83$. Section 2.3 predicts that bill_pass and ss_pass favor the chamber cartel channel and bill_law and ss_law the cross-branch channel; the four differences carry those signs, and none is distinguishable from zero. The cross-branch LES estimate (0.198, standard error 0.059) is the denominator of every cross-branch ratio, so the intervals are wide and the non-rejection is low-powered.

$p = 0.01$) and the joint test rejects; section 6.5 reports the comparison.

The substantive reading is that the outcome profile does not separate a second channel. A genuine cross-branch coordination channel should have moved bills through the second chamber and across the governor’s desk, over-amplifying terminal-law outcomes relative to its own baseline; the profile we observe is the chamber profile, and where the chamber cannot flip, nothing jumps at the cutoff on any outcome. Section 5 provides the second diagnostic: whether the cross-branch pillar responds to chamber-internal cartel-power moderators.

5. Results — Cartel-Power Test and Cross-Pillar Separability

5.1 The chamber cartel premium scales with cartel-power institutional features

The Anzia-Jackman cartel-power index runs from zero (chambers whose standing rules do not codify majority control over hearings, reports, or the floor calendar) to five (chambers whose rules give the majority both committee blocking tools and control of the floor calendar). If the chamber cartel channel of section 2.1 is the right account of why majority-party members are more effective, the size of the chamber-majority premium should scale with this institutional intensity.

The data confirm this prediction in the form the slope can carry it. The index is confined to five realized values — $\{0, 2, 3, 4, 5\}$ — because two of the six summed indicators are disjunctions of the other four, so each tool a chamber adds brings its summary indicator with it. Across these configurations the chamber-majority premium is between four-tenths and two-thirds of a point at the four thinner configurations and 1.12 at the richest, which holds half the analytical sample and twenty-eight of the sixty-seven chambers (Table 4, Panel A). The primary quantitative statement is the continuous slope on the full sample: +0.16 LES points per index unit (standard error 0.02), p below 0.0001 under cluster-robust inference, wild bootstrap interval 0.10 to 0.22. The endpoint contrast summarizes the gradient descriptively, 1.8 times the thinnest configuration’s premium at the richest, but the thinnest cell rests on eight chambers and the endpoint intervals overlap under both Rademacher and Webb six-point weights (Webb, 2023) (Table 4 notes), so the slope is the estimate we rely on. Squire’s professionalism index, taken from the corrected 2015 scores (Squire, 2017) and held fixed within state across the panel, shows a parallel pattern: a top-quartile premium 2.2 times the bottom-quartile premium (Table 4, Panel B).

Parallel gradients on two correlated moderators raise precisely the question the speci-

ficity claim needs answered: is the index capturing cartel rules specifically, or generic chamber capacity? The two moderators are only modestly correlated — a chamber-level correlation of +0.23, with professionalism explaining eight percent of the index’s variance in the estimation sample — and a horse race entering both interactions jointly discriminates them. The majority-premium interaction with the Anzia-Jackman index is +0.23 per standard deviation net of the Squire interaction ($p < 0.001$, wild bootstrap interval 0.13 to 0.32), and the interaction estimated on the Squire-orthogonalized component of the index is identical: the ninety-two percent of the index that professionalism does not explain carries the signal. Squire professionalism also moderates the premium in its own right (+0.10 per standard deviation, $p = 0.001$) net of the index — cartel rules are the dominant institutional feature scaling the premium, and their moderation is not professionalism in disguise. The specificity claim of section 2.3 survives the horse race.

Figure 2 plots both pillars across the five realized configurations: the chamber-majority premium (filled circles) is largest at the richest configuration, while the cross-branch estimates (open circles) scatter around a quarter of a point without a gradient.

A coding-convention check confirms the gradient is not driven by the Anzia-Jackman zero-coding of three chambers (Maine lower, Wisconsin lower, South Carolina upper) for which Anzia and Jackman were unable to code the calendar variables: dropping them entirely leaves the gradient in place (Table 4 note; Appendix H), so the missing-data convention does not produce the result.

The gradient carries the scaling hypothesis from aggregate to member-level, discontinuity-identified data. Anzia and Jackman (2013) built the aggregate version of this test — the same coding of cartel rules, tested on chamber-level majority roll rates; Howard and Provins (2026) link an overlapping set of majority-advantaging chamber rules — four of these tools, folded into a broader nine-rule index — to the effectiveness scores we use, but in cross-sectional ordinary least squares on party-aggregate outcomes; McGrath et al. (2025) use the effectiveness scores with legislative professionalism as a

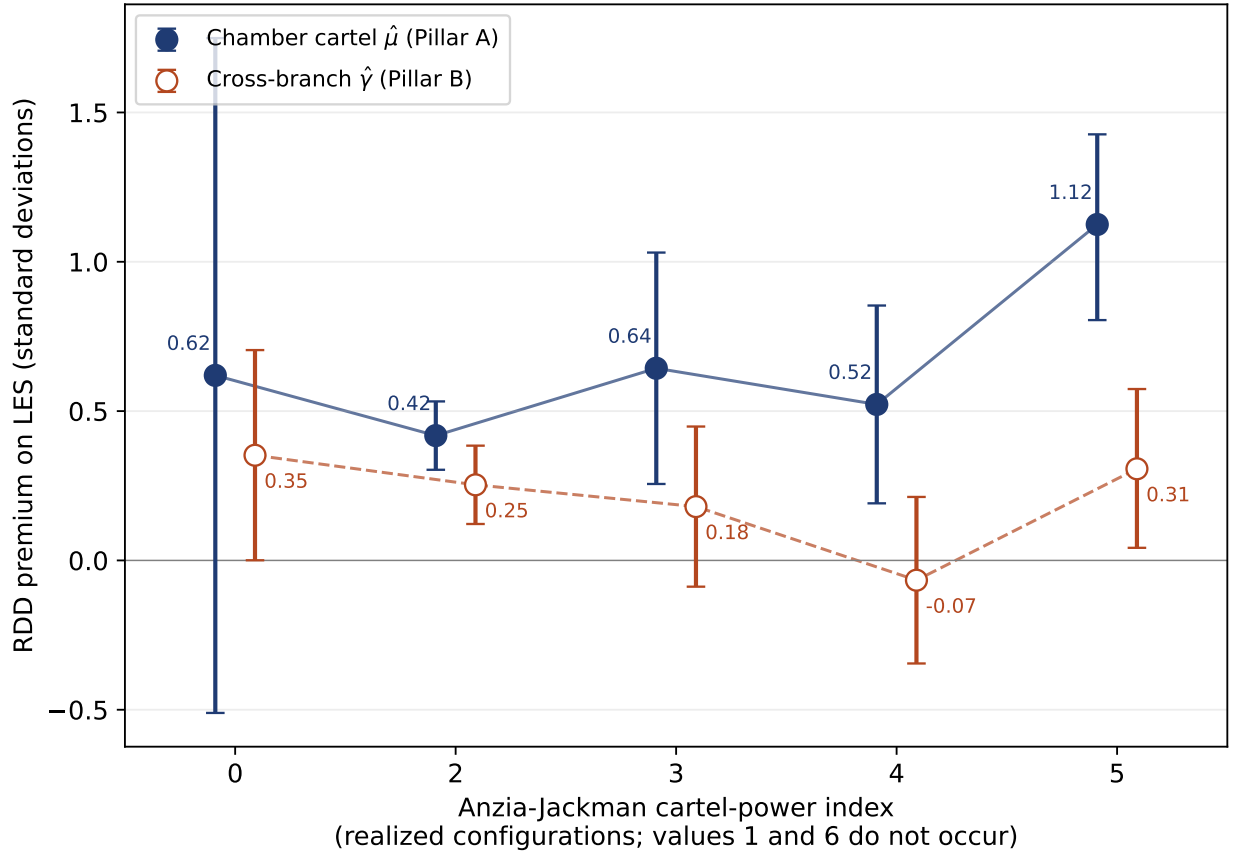


Figure 2: Cross-pillar response to chamber cartel power. RDD estimates at each of the five realized Anzia-Jackman cartel-power configurations $\{0, 2, 3, 4, 5\}$ (values 1 and 6 do not occur), by pillar, plotted on an evenly spaced ordinal axis. Filled circles are the chamber cartel premium ($\hat{\mu}$, Pillar A); open circles are the cross-branch premium ($\hat{\gamma}$, Pillar B). Vertical bars are 95% intervals (wild cluster bootstrap at the two endpoint configurations, analytic ± 1.96 SE in between). The chamber cartel estimates are largest at the richest configuration (endpoint ratio 1.8); the cross-branch estimates show no systematic gradient.

Table 4: Institutional heterogeneity of the two channels

Moderator level	Pillar A ($\hat{\mu}$)	N	G	Pillar B ($\hat{\gamma}$)	N	G
<i>Panel A. Anzia-Jackman cartel-power index (5 realized values)</i>						
0 (cartel-thin)	0.620 (0.160)	2,377	8	0.352 (0.076)	2,552	7
2	0.418 (0.058)	4,143	9	0.253 (0.067)	4,999	6
3	0.643 (0.198)	5,223	15	0.180 (0.137)	7,552	17
4	0.522 (0.169)	2,321	7	-0.066 (0.142)	3,149	8
5 (cartel-rich)	1.125 (0.068)	13,369	28	0.307 (0.076)	14,688	22
Endpoint ratio ($5 \div 0$)	1.82×			0.87×		
Continuous slope (per feature)	+0.161***			+0.035 (n.s.)		
<i>Panel B. Squire professionalism index (quartiles)</i>						
Q1 (least professional)	0.500 (0.080)			0.249 (0.151)		
Q2	0.616 (0.194)			0.491 (0.185)		
Q3	0.937 (0.126)			0.266 (0.106)		
Q4 (most professional)	1.115 (0.108)			0.113 (0.109)		
Q4 $\div$ Q1	2.23×			0.45×		
				(reversed)		

Note: Local-linear RDD estimates within each moderator level at the $\pm 5\%$ primary bandwidth; standard errors in parentheses; N is the cell's member-term count and G its cluster count (state-chamber for Pillar A, state for Pillar B; computed for the Anzia-Jackman panel). The index realizes only five values {0, 2, 3, 4, 5}: two of the six summed indicators are disjunctions of the other four, so an index of 1 is unreachable, and no chamber carries both calendar tools at once, so 6 does not occur. The chamber cartel continuous slope is significant at $p < 0.0001$; the cross-branch slope is indistinguishable from zero ($p = 0.11$). Cells with $G < 12$ have intervals re-estimated with Webb six-point weights (Webb, 2023): the chamber cartel endpoint intervals overlap under either weighting because the index-0 cell rests on eight chambers, no cell's significance changes between weightings, and the cross-branch index-2 cell rests on six clusters with an interval spanning zero, so we do not rely on it. Appendix H reports the interval values, the cell-equality Wald test, and the three-chamber zero-coding check. Stars: *** $p < 0.001$; n.s. not significant.

moderator, but of the gender gap rather than an RDD-identified majority premium. Our contribution is to identify the premium at the individual-member level through a chamber discontinuity and to show that it scales with the cartel-power index while a parallel cross-branch premium does not. The gradient is large by the standards of state-politics heterogeneity findings: the slope implies a rise of about 0.8 points across the index, an amount equal to the chamber-majority premium itself.

5.2 The cross-branch premium shows no cartel-form gradient

The same Anzia-Jackman index is, by construction, an operationalization of chamber-internal cartel power: gatekeeping, hearing denial, reporting denial, and calendar control are all features of how the chamber organizes its own internal agenda. A cross-branch coordination channel does not depend on these features; what it needs is that the member's party simultaneously holds the governorship and the two chambers, institutional alignment across the branches rather than concentration within one. A genuine cross-branch premium should therefore not respond to the index. A state-pillar contrast that carried only the chamber premium would respond, but only in proportion: the chamber gradient scaled by the own-chamber-binding share, about a fifth.

The data show the second pattern. The continuous slope of the cross-branch estimate on the index is +0.04 LES points per feature (standard error 0.02, $p = 0.11$), against +0.16 on the chamber pillar; the chamber slope times the own-chamber-binding share, 0.16 times 0.20, is 0.03, inside the cross-branch slope's interval, and a proportional response of that size cannot be rejected ($p = 0.91$ against the benchmark). Across the five realized configurations the cross-branch estimates show no upward trend: the endpoints differ by less than five-hundredths of a point with overlapping intervals, and the five cells are not jointly distinguishable (cluster-robust Wald $\chi^2(4) = 4.7$, $p = 0.32$; Table 4, Panel A). In the horse race of section 5.1, neither the index (+0.06 per standard deviation, $p = 0.11$) nor professionalism (0.00, $p = 0.94$) moderates the cross-branch estimate, and the Squire quartiles, if anything, reverse.

Read together with section 4, the flat gradient is the third fingerprint of a contrast with no cross-branch component. Both pillars run on the same panel with the same bandwidth filters, and the index is constant within chamber and identical across pillars, so no artifact lets one pillar pick up the gradient and the other miss it; what the state pillar shows is the chamber premium, diluted by the share of cells in which the member's own chamber is the binding institution, and no more. What the evidence cannot do is rule out a cross-

branch gradient smaller than the proportional benchmark, and we do not press the test below that resolution.

6. Robustness

6.1 Aggregation

Each member-term in our primary specification carries one observation regardless of caucus size; an alternative treats each party-caucus within a state-chamber-term as the unit, giving each caucus one vote. The two weightings agree on the effectiveness score (ratios of 0.98 and 1.03) and, on the cross-branch pillar, on the bill stages, but the caucus-mean estimates exceed the member-term primaries by about half on the chamber pillar’s bill outcomes and by roughly a factor of two on the substantive-significance outcomes of both pillars (Table 5, Panel A). A weighted-least-squares check that weights each caucus by its membership reproduces the member-term primary on every outcome within five percent, so the divergence is a weighting finding rather than an implementation issue: the estimators differ only in how caucuses are weighted, so the gap indicates a per-member premium on substantively significant legislation that is larger where caucuses are smaller, consistent with cartel theory’s picture of agenda power concentrated on a productive few rather than spread uniformly.

6.2 Fixed-effect structure

The primary specification for Pillar B uses additive state and year fixed effects: these absorb time-invariant state characteristics and contemporaneous national shocks while preserving the cross-state-year variation that identifies the cross-branch premium. An interactive state-by-year alternative absorbs all state-year-level shocks, leaving identification entirely on within-state-year cross-party asymmetry — non-trivial under the party-

specific construction of section 3, because the two parties' sides of the same state-year carry different running-variable values.

Under the interactive structure the cross-branch estimate rises to 0.84, roughly four times the additive-FE primary and comparable to the chamber cartel premium (Table 5, Panel B). The additive-FE primary compares unified-party-side average effectiveness across state-years; the interactive-FE alternative compares the two parties within a state-year, and in a near-unified state-year that cross-party contrast is a majority-versus-minority comparison in both chambers, so an estimate at the scale of the chamber cartel premium is the signature of a contrast loading on the bundled chamber component (hence the own-chamber-binding subsample's 0.67), not a cross-branch coordination estimate. Among the additive variants, dropping the year effects leaves the estimate unchanged (state-only 0.20) and dropping the state effects lowers it by a fifth (year-only 0.15, no fixed effects 0.16; Appendix B).

6.3 Treatment definition

The most consequential robustness check concerns the party-specific re-orientation of the cross-branch pillar's treatment definition. Under our primary specification (specified in section 3 and motivated in section 2.3), unified government is treated party-specifically: the treatment indicator is one only when the member's own party holds the governorship, the lower chamber, and the upper chamber. Under the Kirkland-Phillips (2020) any-party formulation, unified government is treated as one whenever any party holds all three institutions, regardless of which party the member belongs to.

The two definitions disagree on a small but substantively crucial subset of observations. Of the 19,917 observations classified as control under our primary specification, 1,901 (roughly ten percent) flip from control to treated under the any-party formulation; these are state-years in which the opposite party — not the member's own party — holds unified government. Holding fixed sample, fixed effects, cluster choice, controls, and

bandwidth, the cross-branch estimate falls from the party-specific 0.20 to -0.01 under the any-party alternative (Table 5, Panel C): the any-party contrast is zero. A pairs cluster bootstrap over states rejects equality of the two coefficients (difference 0.21, 95 percent interval 0.13 to 0.35, p below 0.002).

The mechanism clarifies why the re-orientation matters. Under the party-specific construction, an observation in which the opposite party holds unified government is “control”: the member’s team is locked out at every institutional veto point. Under the any-party construction the same observation becomes “treated.” Pooling these substantively opposite cases cancels the contrast: the member whose party gains unified government and the member whose party is locked out by it enter the treated group together, and their gains and losses offset. The cancellation is itself the empirical test justifying the reformulation: the any-party definition is adequate for state-policy outcomes, where what matters is whether some unified coalition exists, but not for individual-member productivity, where what matters is which party holds the coalition. It is also the first of the two steps section 4.1 describes: the party-specific definition makes the member’s own-chamber majority visible at the compound cutoff, and the binding-institution decomposition then shows that it is what the contrast contains.

6.4 Validity battery and identification-axis spectrum

The validity battery summarized at the close of section 3 — covariate continuity on six predetermined characteristics, McCrary density tests at three bandwidths, pre-LES placebo regressions, and a joint 2D-RDD subset analysis — passes five of six and three of six balance cells with the failing cells characterized there, three of three and two of three density windows, the placebo accounted for by status persistence on both pillars, and the joint subset reproducing the chamber premium within seven percent while raising the cross-branch contrast by two-fifths, the direction the compound-cutoff decomposition predicts. The placebo carries a general lesson: panel RDDs with persistent group-level

Table 5: Robustness of the two channels*Panel A. Aggregation — caucus-mean $\div$ member-term estimate*

Outcome	Pillar A	Pillar B
LES	0.98	1.03
bill_pass	1.45	1.05
bill_law	1.50	1.07
ss_pass	1.90	3.54
ss_law	2.00	1.94

Panel B. Fixed-effect structure (Pillar B, $\hat{\gamma}$)

Specification	$\hat{\gamma}$ (SE)	vs. primary
Additive state + year (primary)	0.198 (0.059)	—
State only	0.196 (0.053)	−1%
Year only	0.152 (0.049)	−23%
No fixed effects	0.159 (0.045)	−20%
Interactive state $\times$ year	0.843 (0.304)	+325%

Panel C. Treatment definition (Pillar B, $\hat{\gamma}$)

Definition	$\hat{\gamma}$ (SE)	Wild 95% CI
Party-specific (primary)	0.198 (0.059)	[0.054, 0.333]
Any-party-unified	−0.008 (0.051)	[−0.118, 0.104]

Panel D. Estimator choice (both pillars)

Estimator	Pillar A $\hat{\mu}$ (SE)	Pillar B $\hat{\gamma}$ (SE)
Local-linear, triangular kernel, $\pm 5\%$ (primary)	0.818 (0.108)	0.198 (0.059)
Local-linear, uniform kernel, $\pm 5\%$	0.790 (0.103)	0.180 (0.054)
Local-linear, triangular kernel, $\pm 3\%$	0.821 (0.113)	0.180 (0.057)
Local-constant, uniform kernel, $\pm 5\%$	0.802 (0.063)	0.317 (0.038)
Local-constant, triangular kernel, $\pm 5\%$	0.802 (0.073)	0.285 (0.043)
CCT robust bias-corrected	0.807 (0.029)	0.144 (0.038)

Note: Panel A reports the ratio of the unweighted caucus-mean estimate to the member-term estimate for each outcome and pillar; values near one indicate agreement (section 6.1 discusses the divergence). Panel B: additive fixed-effect variants versus the interactive state $\times$ year specification (section 6.2). Panel C: party-specific versus any-party-unified treatment (section 6.3); a pairs cluster bootstrap rejects equality of the two coefficients (p below 0.002; section 6.3). Panel D rows other than CCT carry the same fixed effects, controls, and clustering as the primaries; the CCT row (robust bias-corrected, triangular kernel, MSE-optimal bandwidths of 7.3 and 4.7 percent) runs without covariates and with nearest-neighbor variance because the mass-point running variable makes the clustered variance matrix singular. The CCT 95% confidence intervals are [0.751, 0.864] for Pillar A and [0.070, 0.218] for Pillar B.

running variables need a covariate discontinuity test on lagged status, a flipped-event subset, or a lagged-status control. Appendix I reports the diagnostics.

The robustness sweep also produces a spectrum of cross-branch estimates: the any-party alternative (section 6.3) is zero, the joint 2D-RDD subset (section 3) is two-fifths above the primary, and the interactive-FE alternative (section 6.2) is four times it. The

compound-cutoff decomposition of section 4.1 governs how the ordering is read: movement up the axis tracks increasing exposure to the compound-margin bundling documented in Table 2, not sharper identification of a coordination premium; the upper rungs are more bundled, not closer to one. Appendix J reports the spectrum.

6.5 Estimator choice

The primary estimator for both pillars is local-linear with a triangular kernel inside the fixed five-percent bandwidth (section 3). Table 5, Panel D reports it alongside the local-constant difference in means within the same window (uniform and triangular kernels), the same local-linear specification with a uniform kernel and at a three-percent bandwidth, and the robust bias-corrected estimator of Calonico et al. (2014) with its own MSE-optimal bandwidth. The chamber cartel premium is estimator-stable: 0.80 local-constant, 0.82 primary, 0.79 uniform-kernel, 0.81 bias-corrected, a range of four percent. The cross-branch contrast is not: 0.32 local-constant, 0.20 primary, 0.18 uniform-kernel, 0.14 bias-corrected. The direction is the one the persistence mechanism of section 3 predicts: a within-window difference in means absorbs the slope of the outcome in the running variable, and on the cross-branch pillar that slope adds six-tenths to the contrast. The estimator matters for one reading in particular: under the local-constant alternative the governor-binding bill_{law} cell of Appendix P is positive (0.71, $p = 0.01$) and the equal-profiles test of section 4.2 rejects, so a local-constant analysis would have supported a terminal-law cross-branch channel that the primary estimator does not. We report both and keep the primary estimator's reading for the reasons given in section 3.

6.6 Forcing-variable construction

Both pillars' running variables are deterministic uniform-shift constructions from observed vote shares, without the Monte Carlo electoral simulation that Kirkland and

Phillips (2020) use. Rebuilding both in the Kirkland-Phillips stochastic style — forty thousand simulated electoral-shock draws, recording for each member the smallest shock that flips the relevant condition — changes neither reading (Appendix O). On the chamber pillar the stochastic estimate is 0.71 against the primary 0.82, thirteen percent below it (0.76 against 0.80 under the local-constant estimator), with overlapping intervals. On the cross-branch pillar the simulated running variable’s sign disagrees with the member’s observed unified status for a fifth of the near-threshold sample, so its stochastic arm is a within-window contrast in observed status rather than an intercept shift at a cutoff; compared with the deterministic construction under the same local-constant estimator, it is 0.31 against 0.32. We retain the deterministic construction for its larger analytical sample and tighter inference.

7. Discussion and Conclusion

The two-pillar design set out to locate the majority advantage and to test whether it scales with the institutional intensity of cartel rules. Three features of the empirical work are worth recording.

First, the chamber cartel premium passes its quantitative test against the Anzia-Jackman cartel-power index. The scaling hypothesis Anzia and Jackman proposed at the chamber level receives here the individual-member, discontinuity-identified test it had not previously had. The gradient is large by the standards of state-politics heterogeneity estimates, robust to the missing-data convention, stable under cluster-robust, wild-bootstrap, and few-cluster Webb inference, and — critically for the cartel interpretation — it survives the professionalism horse race. What is new is not that cartel power operates in the states but that the scaling prediction holds at the level of individual-member productivity; the slope of 0.16 points per codified tool, a rise across the index equal to the chamber premium itself, is the member-level signature.

Second, the advantage is made inside the chamber. The party-specific state design recovers a unified-government contrast of a quarter of the chamber premium, and three independent decompositions agree that the contrast is the chamber premium carried across a compound cutoff: where the governorship is the binding margin, nothing jumps at the cutoff on any outcome; and the cross-branch pillar's outcome profile and its gradient on the cartel-power index are the chamber pillar's, scaled by the own-chamber share. The bound matters as much as the sign: unified control of the remaining branches adds at most an eighth of the chamber premium to a member's effectiveness. The result agrees with the federal description in which the majority advantage is about the same under unified and divided government (Volden and Wiseman, 2014), and identifies it off close margins.

Third, a state-government discontinuity design read at the member level needs two steps the state-policy original does not. Treatment must be party-specific, because the any-party definition pools the member whose party holds unified government with the member locked out by it and cancels the contrast to zero; and the estimate must be decomposed by binding institution, because the party-specific cutoff carries chamber majority for a fifth of the near-threshold sample. The first step moves the contrast from zero to a quarter of the chamber premium and the second moves it back; a member-level study that took either step without the other would report a cross-branch premium that is not there. The same two steps apply to any close-election design in which "party control" bundles a within-chamber and a cross-branch component, in Congress and in comparative presidential systems.

More broadly, the result narrows what electoral victory buys an individual legislator. A narrow chamber majority buys agenda and gatekeeping leverage inside the chamber, scaled by how much of it the chamber's rules codify; unified control of the elected branches buys the governing coalition enactment leverage over policy but no additional effectiveness for the individual member beyond what her chamber majority supplies.

Whether the aggregate output of unified governments differs (section 2.2) is a separate question the member-level null does not settle.

Several caveats apply. The Anzia-Jackman coding of chamber rules is a single cross-section applied to a thirty-eight-year window; rule changes within the window, Colorado's 1988 GAVEL initiative among them, introduce misclassification that attenuates the gradient rather than manufacturing it, and time-invariant miscoding cannot produce a gradient on one pillar and not the other, but the gradient should be read as an average over the window rather than a rule-vintage-specific estimate. Panel seat coverage is below one hundred percent in chambers mixing single- and multimember districts (median sixty-five percent), and LES coverage varies across states; the gap follows district structure, not member productivity, and the chamber premium rises rather than falls as thinly covered chambers are excluded (Appendix L). The forcing variables are deterministic uniform-shift constructions; section 6.6 shows that rebuilding them in the Kirkland-Phillips Monte Carlo style leaves both readings in place while preserving the close-margin as-if-random property the design relies on (Lee, 2008). Term-limited states show some mechanical compression in late-career terms that our estimates pool with non-term-limited states; the term-limit split shows no significant premium difference (Appendix H).

The two-pillar architecture suggests two extensions. The first is to bring the decomposed cross-branch design to other member-level outcomes, where a coordination premium may exist even though effectiveness shows none: committee-chair appointment, roll-call cohesion with leadership, and distributive returns to the member's district. The second is to bring the chamber cartel test to other settings where chamber rules vary cross-sectionally and contemporaneously — measures capturing committee assignment power, conference-committee membership, and presiding-officer authority would test the prediction against richer variation than the Anzia-Jackman coding exploits. Both extensions share the design's structural requirement — close electoral margins on chamber compo-

sition plus a separate close-margin forcing variable for the cross-branch institutional unit — which generalizes beyond U.S. state legislatures, although member-level effectiveness panels matched to electoral-margin panels are uncommon elsewhere.

References

- Aldrich, J. H. (1995). *Why Parties? The Origin and Transformation of Political Parties in America*. University of Chicago Press, Chicago.
- Alpino, M. and Crispino, M. (2024). The role of majority status in close election studies. *Political Analysis*, 32(1):148–155.
- Anzia, S. F. and Jackman, M. C. (2013). Legislative organization and the second face of power: Evidence from US state legislatures. *Journal of Politics*, 75(1):210–224.
- Bucchianeri, P., Volden, C., and Wiseman, A. E. (2025). Legislative effectiveness in the american states. *American Political Science Review*, 119(1):21–39.
- Calonico, S., Cattaneo, M. D., and Titiunik, R. (2014). Robust nonparametric confidence intervals for regression-discontinuity designs. *Econometrica*, 82(6):2295–2326.
- Cameron, A. C., Gelbach, J. B., and Miller, D. L. (2008). Bootstrap-based improvements for inference with clustered errors. *Review of Economics and Statistics*, 90(3):414–427.
- Chin, M. J. and Shi, L. (2025). Average and heterogeneous effects of political party on state education finance and outcomes: Regression discontinuity evidence across U.S. election cycles. *Economics of Education Review*, 105:102636.
- Coleman, J. J. (1999). Unified government, divided government, and party responsiveness. *American Political Science Review*, 93(4):821–835.
- Cox, G. W., Kousser, T., and McCubbins, M. D. (2010). Party power or preferences? quasi-experimental evidence from american state legislatures. *Journal of Politics*, 72(3):799–811.
- Cox, G. W. and McCubbins, M. D. (1993). *Legislative Leviathan: Party Government in the House*. University of California Press, Berkeley.

- Cox, G. W. and McCubbins, M. D. (2005). *Setting the Agenda: Responsible Party Government in the U.S. House of Representatives*. Cambridge University Press, Cambridge.
- de Magalhães, L. (2021). Estimating slim-majority effects in U.S. state legislatures with a regression discontinuity design under local randomization assumptions. *Political Science Research and Methods*, 9(3):665–674.
- Hitt, M. P., Volden, C., and Wiseman, A. E. (2017). Spatial models of legislative effectiveness. *American Journal of Political Science*, 61(3):575–590.
- Hopkins, D. J. (2018). *The Increasingly United States: How and Why American Political Behavior Nationalized*. University of Chicago Press, Chicago.
- Howard, N. and Provins, T. (2026). When process becomes power: Rules, parties, and legislative effectiveness. Working paper, Center for Effective Lawmaking. Available at <https://thelawmakers.org/legislative-research/when-process-becomes-power-rules-parties-and-legislative-effectiveness>.
- Kirkland, P. A. and Phillips, J. H. (2020). A regression discontinuity design for studying divided government. *State Politics & Policy Quarterly*, 20(3):356–389.
- Lee, D. S. (2008). Randomized experiments from non-random selection in U.S. house elections. *Journal of Econometrics*, 142(2):675–697.
- McCrary, J. (2008). Manipulation of the running variable in the regression discontinuity design: A density test. *Journal of Econometrics*, 142(2):698–714.
- McGrath, R. J. and Ryan, J. M. (2019). Party effects in state legislative committees. *Legislative Studies Quarterly*, 44(4):553–592.
- McGrath, R. J., Ryan, J. M., and Wrighten, J. (2025). High hurdles: Legislative professionalism and the effectiveness of women state legislators. *Journal of Politics*. Forthcoming.

- Napolio, N. G. and Grose, C. R. (2022). Crossing over: Majority party control affects legislator behavior and the agenda. *American Political Science Review*, 116(1):359–366.
- Repetto, L. and Sosa Andrés, M. (2023). Divided government, polarization, and policy: Regression-discontinuity evidence from US states. *European Journal of Political Economy*, 80:102473.
- Rohde, D. W. (1991). *Parties and Leaders in the Postreform House*. University of Chicago Press, Chicago.
- Squire, P. (2017). A squire index update. *State Politics & Policy Quarterly*, 17(4):361–371.
- Volden, C. and Wiseman, A. E. (2014). *Legislative Effectiveness in the United States Congress: The Lawmakers*. Cambridge University Press, Cambridge.
- Webb, M. D. (2023). Reworking wild bootstrap-based inference for clustered errors. *Canadian Journal of Economics*, 56(3):839–858.

This appendix reports the full grid of robustness checks summarized in section 6 of the main text, along with the validity-battery diagnostics summarized at the close of section 3. Each subsection corresponds to one dimension of the identification-sensitivity sweep specified in advance in our analysis protocol. Two disclosures belong at the top. First, the estimator. The specification fixed before estimation was a local-constant difference in means within the fixed five-percent window; it was changed to the local-linear specification of section 3 before the re-estimation reported in this version, on the grounds given there (standard practice, the original state-government design’s own local-linear estimation, and the persistence-driven boundary bias), at a point when the local-linear headline estimates were known and the binding-institution decomposition of Table 2 under that estimator was not. The local-constant estimates are reported beside the primary ones in section 6.5 and, where the reading differs, in section P. Second, the data. The effectiveness scores are the 2026 per-state release of the SLES, which extends the panel through the 2023–2024 biennium, adds Kansas, and revises scores in the earlier window for New York (a companion-bill scoring rule), Washington, Colorado, and Alaska; section L records the construction.

A Bandwidth sensitivity

We report the chamber cartel premium and the cross-branch coordination premium across four primary bandwidths: $\pm 1\%$, $\pm 3\%$, $\pm 5\%$ (primary), and $\pm 10\%$ of statewide vote share. The primary bandwidth is specified in section 3.

For the chamber cartel pillar, the estimate ranges from 0.7726 at $\pm 1\%$ to 0.8354 at $\pm 10\%$, an overall spread of 7.7%: the narrow window gives the smallest estimate (-5.5% versus primary) and the broad window the largest ($+2.2\%$). All four bandwidth estimates pass the pre-specified $\pm 20\%$ tolerance.

For the cross-branch pillar the estimate is bandwidth-sensitive: 0.1984 at the primary

5%, 0.1796 at $\pm 3\%$ (-9.5%), 0.2091 at $\pm 10\%$ ($+5.4\%$), and -0.0059 at $\pm 1\%$, the one bandwidth that fails the tolerance. Section 4.1 explains why a bundled contrast is fragile in a narrow window: what the contrast contains depends on which institution is binding in the state-years the window happens to retain, and the $\pm 1\%$ window retains few. We report the primary 5% as the headline and the $\pm 1\%$ cell as a caveat; under the decomposition of section P the clean cross-branch component is zero at every bandwidth this grid uses.

B Fixed-effect structure

For each pillar we report the estimate under five fixed-effect specifications: the additive primary, an interactive alternative, two reduced-FE alternatives, and a no-FE specification. The chamber cartel pillar’s estimate ranges from 0.7875 to 0.8551 across the five variants, well within the pre-specified $\pm 30\%$ tolerance for additive FE and $\pm 50\%$ for interactive FE. An exploratory specification grid run before the main analysis had documented the state-chamber-by-year interactive FE as a parallel sensitivity yielding 0.8551; the value here reproduces that result within rounding.

For the cross-branch pillar, the four additive variants — the state-plus-year primary, state-only, year-only, and no-FE — all produce estimates in the 0.1519 to 0.1984 range, within roughly $\pm 23.4\%$ of the primary. The interactive state-by-year FE produces 0.8431, a $+324.9\%$ deviation discussed at length in section 6.2 of the main text. The interactive estimate is reported as a parallel sub-question identifying off within-state-year asymmetric-unified variation rather than as a falsification of the primary.

C Cluster sensitivity

We report three inference methods on every main result per our pre-specified inference protocol. The chamber cartel pillar uses state-chamber as the primary cluster ($G = 67$), state as the secondary ($G = 38$), and wild cluster bootstrap with 999 Rademacher replicates and restricted residuals as the third. The point estimate is identical to four decimals across the three methods; only standard errors and inference machinery differ. The wild bootstrap symmetric percentile-t 95% confidence interval is [0.458, 1.192].

The cross-branch pillar uses state as the primary cluster ($G = 41$ in the analytical sample; forty-two states in the population panel), state-year as the secondary ($G = 306$), and wild cluster bootstrap with 999 Rademacher replicates as the third. The point estimate is identical to four decimals across the three methods; the wild bootstrap 95% interval is [0.054, 0.333]. Because state cluster count is moderate, the wild bootstrap is not optional for the cross-branch pillar; it is reported as the canonical inference robustness alongside the standard CR1 results.

D Donut radius

A donut RDD re-estimates the design after excluding observations immediately adjacent to the threshold, on the rationale that mechanical features of the running variable in a small radius might violate the as-if-random property. We test two donut radii (0.5% and 1.0%) for each pillar.

For the chamber cartel pillar, a pre-specified concern was that tied-chamber observations sitting exactly at the threshold might bias the estimate even though they are excluded from the primary specification. The donut analysis at both radii produces estimates of 0.8436 and 0.7806 — both within 4.5% of the primary 0.8177. The tied-chamber boundary is boundary-insensitive at the primary cutoff.

For the cross-branch pillar, a parallel pre-specified check tests boundary sensitivity at

the same two radii. The donut produces estimates of 0.2314 and 0.1952, both within 16.6% of the primary 0.1984. The cross-branch boundary is similarly boundary-insensitive.

E Forcing-variable default-choice sensitivities

The Pillar B forcing-variable construction involves four default choices that were specified in advance; all four are tested. The fourth, treatment timing, aligns treatment to the January session start; aligning it instead to inauguration dates by dropping the four off-cycle states moves the estimate by less than one percent (0.1978 against 0.1984).

The governor roll-forward check tests whether the roll-forward imputation of the sitting governor in off-election state-years affects the cross-branch estimate. Dropping all off-election state-years produces an estimate of 0.2637 (+32.9% versus primary, wild $p = 0.001$), outside the pre-specified $\pm 20\%$ tolerance. The imputation therefore matters for the magnitude of the bundled contrast, and we record the check as a caveat rather than a validation; the primary estimate keeps the off-election state-years because dropping them removes most of the governor-binding cells on which section P's clean estimate rests.

The seven-state placebo tests whether the seven zero-coverage states' exclusion produces a meaningful sample bias. We run a one-hundred-replicate random-seven-state-drop bootstrap on the forty-one-state primary set. The mean estimate across the one hundred replicates is 0.2028 (standard deviation 0.0327), placing the actual primary (0.1984) at $z = -0.13$ in the placebo distribution. The seven excluded states behave statistically like a random seven-state loss; the exclusion is benign.

The any-party-unified treatment definition is reported in the main text as section 6.3.

F McCrary density tests

We run density-continuity checks on the running variables at the threshold ($Z = 0$) for both pillars, in the spirit of the manipulation test of McCrary (2008), at three windows: $\pm 0.5\%$, $\pm 1\%$, and $\pm 2\%$. The statistic is the log ratio of observation counts on either side of the threshold, judged against a 95% Poisson interval; it is a count-symmetry diagnostic rather than the local-linear density-discontinuity estimator of that paper.

For the chamber cartel pillar, all three windows pass: log-density ratios are -0.0748 , -0.0046 , and $+0.0444$ respectively, with each in the 95% Poisson confidence interval of zero. There is no evidence of manipulation of the chamber cartel forcing variable at the threshold.

For the cross-branch pillar, two of three windows pass: log-density ratios are $+0.0456$ at $\pm 0.5\%$ (pass), -0.2014 at $\pm 1\%$ (fail; confidence interval $[-0.2486, -0.1541]$ excludes zero), and -0.0069 at $\pm 2\%$ (pass). The $\pm 1\%$ asymmetry reflects the small-bin sample composition ($n_L = 3,827$, $n_R = 3,129$) on the party-specific running variable; both adjacent windows pass cleanly, and the primary 5% bandwidth identification is not threatened by the narrow-window asymmetry. We report the diagnostic as a minor caveat rather than a substantive identification concern.

G Joint 2D-RDD subset

A subset of state-chamber-terms satisfies both close-margin conditions simultaneously: $|Z_A| < 5\%$ on the chamber cartel forcing variable and $|Z_B| < 5\%$ on the cross-branch forcing variable. This joint subset contains 15,542 member-terms across thirty-two states and provides a within-sample test of whether the two channels identify in the same observations.

In this joint subset, the chamber cartel estimate is 0.8749 ($+7.0\%$ versus the full-sample primary, wild 95% confidence interval $[0.480, 1.290]$). The cross-branch estimate is 0.2838

(+43.1% versus the full-sample primary, wild 95% confidence interval [0.033, 0.559]), outside the $\pm 30\%$ tolerance. Both pillars identify in the same observations; the cross-branch estimate’s upward movement in the joint subset is the one the compound-margin bundling of main-text Table 2 predicts, because conditioning on the chamber margin being close raises the share of near-threshold observations whose binding institution is a chamber, and it extends the identification-axis spectrum of section J. It should not be read as a cleaner cross-branch estimate.

H Heterogeneity robustness: zero-coding, cell-level inference, and the remaining pre-specified moderators

The Anzia-Jackman cartel-power index is constructed as the sum of six binary indicators. Three chambers in our panel — Maine’s lower chamber, Wisconsin’s lower chamber, and South Carolina’s upper chamber — have missing data on three of the six indicators because their standing rules do not codify calendar-control tools in the form the index expects. The convention is to code missingness as zero, treating absence of codified power as the institutionally relevant feature.

We test whether the chamber cartel heterogeneity gradient depends on this convention by dropping the three chambers entirely from the heterogeneity analysis. Under the full sample (F1), the low-to-high cartel-power gradient is +0.534; under the three-chamber-drop subset (F2), the gradient is +0.472 (the low cell moves from 0.52 to 0.58, the high cell is unchanged at 1.05). The ratio $F2/F1 = 0.88$: the gradient is slightly weaker after dropping the chambers, so the missing-data convention does not produce the result; the three chambers contribute modestly to it rather than against it.

Cell-level inference detail for main-text Table 4. The chamber cartel endpoint intervals are, at index 0, Rademacher [−0.511, 1.749] and Webb [−0.458, 1.571], and at index 5, [0.805, 1.427] and Webb [0.828, 1.405]: they overlap under either weighting, because the

index-0 cell rests on eight chambers, and no cell's significance changes between weightings. The cross-branch endpoint intervals overlap. The cross-branch index-2 cell (0.253) rests on six clusters with an interval spanning zero under both weightings (Rademacher $[-0.035, 0.542]$; Webb $[-0.045, 0.549]$). The five cross-branch cells are not jointly distinguishable: cluster-robust Wald $\chi^2(4) = 4.68$, $p = 0.322$. The proportional benchmark of section 5.2, the chamber slope times the own-chamber-binding share, is +0.039; the cross-branch slope of +0.035 does not differ from it ($t = 0.11$, $p = 0.91$).

Our analysis plan pre-specified four heterogeneity moderators: the Anzia-Jackman cartel-power index, Squire professionalism, term-limit status, and chamber type. Main-text Table 4 reports the first two. The remaining two are reported here, re-estimated at the primary specification so that the cells are directly comparable with that table rather than with the exploratory grid in which they were first run. Term-limit status is the static fifteen-state classification specified in our analysis plan; like the Anzia-Jackman coding it is a single cross-section applied to the full window, so it misclassifies the early years of states that adopted limits during the 1990s and the later years of the states whose limits were repealed or judicially invalidated.

Cell	Chamber cartel (μ)			Cross-branch (γ)		
	Estimate	N	G	Estimate	N	G
Non-term-limited states	0.892	18,876	44	0.241	23,968	29
Term-limited states	0.642	8,557	23	0.095	8,972	12
<i>Interaction</i> ($\times$ term-limited)	-0.140	$p = 0.30$		-0.026	$p = 0.72$	
Lower chambers	0.793	20,020	32	0.154	24,678	41
Upper chambers	0.920	7,413	35	0.355	8,262	41
<i>Interaction</i> ($\times$ upper)	-0.005	$p = 0.97$		$+0.068$	$p = 0.11$	

Neither moderator produces a detectable gradient on either pillar. On the chamber cartel pillar the premium is 0.892 in states without term limits against 0.642 in states with

them, a difference of -0.140 that is not distinguishable from zero ($p = 0.30$; wild bootstrap $p = 0.41$), and 0.793 in lower chambers against 0.920 in upper chambers, a raw gap that the pooled interaction, once the fixed effects are in, places at -0.005 ($p = 0.97$). On the cross-branch pillar the corresponding contrasts are 0.241 against 0.095 (difference -0.026 , $p = 0.72$) and 0.154 against 0.355 (difference $+0.068$, $p = 0.11$). None of the four interactions separates from noise at this sample size and we build nothing on them. The rows complete the pre-specified sweep; they do not add a finding.

I Pre-LES placebo and persistence verification (both pillars)

A standard pre-LES placebo regresses the previous term's effectiveness score on the current treatment indicator; under the null, current treatment cannot predict a prior-period outcome. On the primary local-linear specification the placebo is $+0.279$ on the chamber pillar (standard error 0.101 , $p = 0.006$; $N = 19,962$ member-terms with a lagged score, 66 clusters) and $+0.120$ on the cross-branch pillar (standard error 0.084 , $p = 0.15$; $N = 24,402$, 41 clusters). Read literally, the chamber pillar fails.

Both coefficients are the size that status persistence predicts. Majority control persists from one close-margin term to the next for 72 percent of members on the chamber pillar and unified status for 75 percent on the cross-branch pillar, so members treated today were disproportionately treated last term, and the placebo picks up last term's premium. The direct test is a covariate discontinuity in lagged status itself (section P, part 3). At the chamber cutoff the lagged-majority indicator jumps by $+0.336$ (standard error 0.131 , $p = 0.01$; wild interval 0.028 to 0.634), and the primary premium times that jump, $0.818 \times 0.336 = 0.275$, matches the measured placebo of 0.279 . At the state cutoff the lagged-unified indicator jumps by $+0.179$ (standard error 0.107 , $p = 0.10$); the predicted placebo is 0.035 and the measured 0.120 sits inside its own interval around zero, so the arithmetic

is looser there but the placebo itself is not distinguishable from zero.

Three further diagnostics agree. First, controlling for lagged status, the current-treatment coefficient in the placebo regression is +0.023 on the chamber pillar ($p = 0.67$) and +0.065 on the cross-branch pillar ($p = 0.40$), while the lagged-status coefficient is +0.77 and +0.34, the premia themselves: prior-term status on prior-term effectiveness is what the placebo was measuring. Second, in chamber-terms where control actually flipped between $t-1$ and t ($N = 5,622$ member-terms on the chamber pillar; 5,980 on the cross-branch pillar), the placebo is -0.871 (standard error 0.115) and -0.227 (standard error 0.138, $p = 0.10$), the mirror images a genuine premium must produce when prior status was opposite and anticipation would not. Third, the primaries are unchanged by the imbalance: with lagged status as a control the chamber premium is 0.879 (lagged-status coefficient -0.042 , standard error 0.030) and the cross-branch contrast 0.231 (-0.014 , standard error 0.029), so a lagged-status imbalance of 0.34 moves the primary by about 0.01; and the current premium in flipped chamber-terms is not below that in retained ones (1.03 against 0.83 on the chamber pillar; 0.31 against 0.16 on the cross-branch pillar). Controlling instead for the lagged score (0.725 and 0.145) over-adjusts, because the lagged-score imbalance is lagged status's effect and the regression attributes it to persistent ability; we report it for completeness and do not rely on it.

The standard pre-LES placebo is mis-specified for designs whose running variable is a chamber- or state-level status with high period-to-period persistence: it detects persistence even when the flip itself is as-if-random. Panel RDDs of this kind should report the covariate discontinuity in lagged status, a flipped-event subset, or a lagged-status control. The identifying assumption is supported on both pillars.

J Cross-branch identification-axis spectrum

The robustness sweep produces four cross-branch estimates that order along a single identification axis. The any-party-unified alternative of section 6.3 produces -0.0077 — the loosened-treatment-definition rung. The cross-state-year primary of section 4.1 produces 0.1984 — the headline. The joint 2D-RDD subset of section G above produces 0.2838 — the rung that conditions on close margins for both pillars. The interactive state-by-year fixed-effect alternative of section 6.2 produces 0.8431 — the rung that identifies only off within-state-year asymmetry.

The compound-cutoff decomposition of main-text Table 2 governs how this ordering should be read. Movement up the axis tracks increasing exposure to compound-margin bundling, not increasingly clean identification of a coordination premium: the joint subset mechanically raises the share of near-threshold observations whose binding institution is a chamber (section G above), and the interactive-FE contrast reaches the scale of the chamber cartel premium itself — by the decomposition’s own logic, the signature of a chamber-loaded contrast rather than a sharper cross-branch estimate. Read against that decomposition, the axis is anchored at zero: the any-party rung coincides with the governor-binding estimate of main-text Table 2 (-0.075) and with the other-chamber-binding estimate (-0.051), the three designs that remove the member’s own chamber majority from the contrast, and every rung above them adds bundled chamber majority rather than coordination. The primary remains the reported contrast because it respects the party-specific construction motivated in section 2.3 while staying closest to the cross-state-year variation that the additive fixed-effect structure preserves; its composition is what Table 2 and section P make explicit.

K Specification table summary

The complete robustness sweep is summarized as Pass/Fail or interpretable-divergence verdicts in the table below.

Dimension	Chamber cartel (μ)	Cross-branch (γ)
Bandwidth $\pm 1\%$ / $\pm 3\%$ / $\pm 5\%$ / $\pm 10\%$	4/4 PASS	3/4 PASS ($\pm 1\%$ FAIL, section A)
Fixed-effect structure (5 variants)	5/5 PASS	4/5 PASS (interactive +324.9% reported as parallel sub-question, section 6.2)
Cluster (primary / secondary / wild)	Invariant; wild 95% CI [0.458, 1.192]	Invariant; wild 95% CI [0.054, 0.333]
Donut radius (0.5% / 1%)	PASS (within 4.5% of primary)	PASS (within 16.6% of primary)
Governor roll-forward (cross-branch only)	—	FAIL (+32.9%; caveat, section E)
Seven-state placebo (cross-branch only)	—	PASS ($z = -0.13$ against placebo distribution)
Any-party-unified treatment (cross-branch only)	—	Substantive finding — section 6.3
Covariate continuity (six covariates per pillar)	5/6 PASS (appointment status fails; section M)	3/6 PASS (Squire professionalism fails on magnitude, gender and chamber size on p; section M)

Dimension	Chamber cartel (μ)	Cross-branch (γ)
McCrary density (three bin widths)	3/3 PASS	2/3 PASS ($\pm 1\%$ asymmetry caveat)
Pre-LES placebo	Persistence, verified by the lagged-status discontinuity (section I)	PASS; same mechanism (section I)
Joint 2D-RDD subset (N = 15,542)	PASS (+7.0%)	+43.1% (outside $\pm 30\%$; compound-margin bundled rung, section J)
Aggregation (caucus-mean vs member-term)	LES within 2%; bill stages 45–50% above; ss outcomes about double (membership-weighted WLS reproduces every cell within 5%; section 6.1)	LES and bill stages within 7%; ss outcomes double or more (WLS within 5%)
Three-chamber-NaN heterogeneity (cartel pillar only)	F2/F1 = 0.88 (gradient slightly weaker after drop)	—
Chamber-term sorting (three bins around the flip margin)	1/3 PASS (section M)	—
Remaining pre-specified moderators (term limits, chamber type)	No detectable gradient (section H above)	No detectable gradient (section H above)

Dimension	Chamber cartel (μ)	Cross-branch (γ)
Anzia-Jackman coding vintage (window split at 2003)	Slope +0.132 early / +0.185 late; era difference $p = 0.11$ (section N)	Flat in both halves (section N)
Operating-majority flag vs seat counts	drop +7.8% (PASS); seat re-code -20.5% (outside $\pm 20\%$, wild interval contains the primary; section L)	PASS (drop +3.9%, section L)

The chamber cartel premium is stable through the complete sweep, with the seat-count re-code as the one variant that moves it beyond the tolerance, and the cross-branch contrast is stable through inference, donut, placebo and the seven-state placebo. The cross-branch pillar’s divergences ($\pm 1\%$ bandwidth, governor roll-forward, joint subset, interactive FE, any-party treatment) are all movements in a bundled contrast and are read through the compound-cutoff decomposition of section J and section P: each moves the share of near-threshold cells in which the member’s own chamber is the binding institution, and none is a movement in a coordination premium, which section P places at zero on every outcome.

L Sample construction details

The analytical panel is built from three primary sources that we splice into a single member-term file. The Bucchianeri-Volden-Wiseman SLES dataset, in its 2026 per-state release, provides the State Legislative Effectiveness Score and the four funnel-component outcomes (bill_pass, bill_law, ss_pass, ss_law) for each member-term from the 1987–1988 through the 2023–2024 biennium, alongside the legislator-level identifiers, party affilia-

tion, chamber, and biennium; the build merges the per-state score files with the release's covariate files and records the source of every filled field. The Klarner State Legislative Election Returns database provides the district-level vote shares from which both forcing variables are constructed. The Kirkland-Phillips replication archive provides the gubernatorial election outcomes through 2012, which we splice with Dave Leip's Atlas data for 2013-2022; winner agreement on the overlap window is one hundred percent.

The chamber cartel forcing variable (Pillar A) is constructed as the smallest uniform statewide vote-share shift that would flip chamber majority control. For each chamber-term, we compute the absolute value of this shift and assign it a sign according to the member's party: positive when the member's party currently holds the chamber majority, negative when she does not. The construction is closed-form per state-chamber-term and requires no Monte Carlo simulation.

The cross-branch forcing variable (Pillar B) is constructed as the smallest uniform statewide shift that would flip the member's party-specific unified-government status, computed from the three institution-specific margins: the shift that would flip the governorship, the shift that would flip lower-chamber majority, and the shift that would flip upper-chamber majority. The construction is asymmetric across the cutoff because the flip event itself is asymmetric. For a member whose party currently holds unified government, losing any one institution breaks unified status, so the flip distance is the minimum across the three margins. For a member whose party does not hold unified government, attaining it requires flipping every institution her party does not hold, so the flip distance is the maximum across the unheld margins. The sign is positive when the member's party currently holds unified state government and negative when she does not. The binding margin — the institution attaining the minimum on the unified side, or the maximum over unheld institutions on the non-unified side — can shift across institutions year by year; this is by design, because the substantive question is how close the member's-party-unified configuration is to flipping, not which particular institution

is closest to flipping. Table 2 in the main text discloses the binding-institution composition of the near-threshold sample and the compound-margin consequence: for the own-chamber-binding subset, crossing the cutoff flips chamber majority together with unified status.

Seven states (Arizona, Idaho, North Dakota, New Jersey, South Dakota, Washington, West Virginia) drop out of the Pillar B analytical sample because their lower- or upper-chamber composition cannot be parsed into smallest-shift-to-flip terms under the Klarner database's structure: these chambers either use mixed-MMD seat structures or jungle-primary elections that the database does not pre-process. These seven states drop out of Pillar A's chamber RDD subsample for the same reason, propagating through the min operator to Pillar B. The exclusion is documented in our analysis plan and tested with the placebo distribution in section E above; the seven excluded states behave statistically like a random seven-state loss from the full panel.

Tied chambers (where the two parties hold equal seat counts) are excluded from the primary specification because they reflect coalition arrangements rather than as-if-random close-margin assignment. The exclusion is specified in advance; donut RDD checks at 0.5% and 1% radii (section D above) confirm that the exclusion line is boundary-insensitive at the five-percent primary bandwidth.

Chamber majority is assigned from the SLES operating-majority flag wherever the dataset records one, falling back to the seat-derived majority otherwise. The two criteria disagree on forty-nine of the panel's 1,206 chamber-terms. Twenty are exact seat ties, for which no seat-derived winner exists; ten are chamber-terms for which SLES records no operating-majority flag at all; the remaining nineteen are reversals in which the party that organized the chamber is not the party holding the most seats. Those nineteen are the power-sharing and cross-party coalition arrangements the flag exists to capture, the best known in our window being the New York Senate terms organized by a Republican-Independent Democratic Conference coalition and the 2013-2014 Washington Senate Ma-

jority Coalition Caucus. Two features of the comparison matter for how the disagreement should be read. First, the seat counts on the other side of it tally every member who served during a term, including mid-term replacements, and so exceed nominal chamber size by a median of 1.7 percent; a one-seat margin can invert on that basis alone, which makes the flag not merely a different convention but the more reliable record of who held the chamber. Second, disputed chamber-terms are close-margin by construction and are therefore over-represented near the threshold: 1,510 of the 27,433 chamber cartel observations (5.5 percent, from twenty-two chamber-terms) and 1,314 of the 32,940 cross-branch observations (4.0 percent, from thirty-five chamber-terms) come from one. The twenty seat-tie cases carry a flip distance of exactly zero and are already removed by the tied-observation rule above.

Two variants bound what the choice of criterion can do to the estimates. Dropping every observation belonging to a disputed chamber-term leaves the chamber cartel premium at 0.882 ($N = 25,923$, +7.8 percent against the primary) and the cross-branch premium at 0.206 ($N = 31,626$, +3.9 percent). Re-coding majority status from seat counts instead of the flag — which reverses the treatment label, and with it the sign of the running variable, for 1,289 of the 27,433 chamber cartel observations — yields 0.650, 20.5 percent below the primary, outside the $\pm 20\%$ tolerance the other grids apply, with a wild bootstrap interval of $[0.214, 1.032]$ that contains the primary estimate; the re-code is the more consequential variant because it reverses treatment in exactly the closest chamber-terms, where the flag and the seat count disagree and where the local-linear fit has the least support. We do not re-code the cross-branch pillar: its treatment and its party-specific forcing variable are derived jointly from chamber control, so re-coding would mean rebuilding the running variable rather than relabelling a column, and the drop variant bounds that channel. At the member level, a legislator's own SLES majority flag disagrees with the chamber-level assignment in 999 of 94,107 member-terms, or 1.1 percent.

The following cascade traces each pillar's analytic sample from the raw member-term

panel. The bandwidth restriction is the design’s identifying device, not an incidental data loss: it is by design that the $\pm 5\%$ window retains the close-margin observations on which the discontinuity is identified. The forcing-variable step is where the seven zero-coverage states and the mixed-member-district chambers drop out, for the coverage reasons described above.

Construction step	Pillar A (N)	Pillar B (N)
Raw member-term panel	94,107	94,107
Valid LES, two-party, treatment	93,800	87,549
Window 1987–2024	93,800	87,549
Forcing variable defined	76,190	78,075
Within $\pm 5\%$ bandwidth	28,643	33,583
Tied excluded (analytic sample)	27,433	32,940

The two pillars draw on the same raw panel but reach slightly different analytic samples: Pillar B retains more observations at the bandwidth (32,940 versus 27,433) because the party-specific unified-government margin places more member-terms within five percent of a flip than the chamber-majority margin does. Seat coverage, the share of a chamber’s seats for which the Klarner database yields a usable district margin, has a median of 0.65 across chamber-terms; restricting the chamber cartel pillar to chambers with coverage of at least 25, 50 and 75 percent gives 0.853 (N = 26,420; G = 65), 0.903 (N = 23,147; G = 58) and 0.934 (N = 11,416; G = 40) against the primary 0.818, so thin coverage, if anything, understates the premium.

M Covariate balance and chamber-term sorting

The validity battery summarized at the close of section 3 reports covariate continuity as five of six passes on the chamber cartel pillar and three of six on the cross-branch pillar. The full table is below. Each predetermined covariate is regressed on treatment at the

pillar's primary specification, with no controls, since the covariate is itself the outcome. A cell passes when the treatment coefficient is statistically indistinguishable from zero ($p > 0.05$) and substantively small (absolute coefficient below one tenth of an outcome standard deviation); both conditions were fixed in advance. For the discrete-seat-count margins that arise in some chambers, de Magalhães (2021) adapts the local-randomization framework of Cattaneo, Frandsen, and Titiunik to legislative seat counts: assignment near the threshold may be treated as if random within a window whose predetermined covariates balance, a condition he finds satisfied in state houses but rejected in state senates. The table below is the analogous check on our own design.

Covariate	b	SE	p	b /SD	Verdict
<i>Panel A. Chamber cartel pillar (N = 27,433; G = 67)</i>					
Squire professionalism	−0.0032	0.0023	0.167	0.021	Pass
Seniority (terms served)	0.1130	0.1302	0.385	0.034	Pass
Female	−0.0088	0.0177	0.619	0.020	Pass
Ideology (Shor-McCarty)	0.0765	0.0422	0.070	0.080	Pass
Chamber size	0.0009	0.0014	0.518	0.000	Pass
Appointed member	−0.0148	0.0065	0.023	0.104	Fails on p and magnitude
<i>Panel B. Cross-branch pillar (N = 32,940; G = 41)</i>					
Squire professionalism	−0.0236	0.0185	0.201	0.164	Fails on magnitude
Seniority (terms served)	0.0068	0.1426	0.962	0.002	Pass
Female	−0.0372	0.0168	0.026	0.087	Fails on p
Ideology (Shor-McCarty)	0.0823	0.0436	0.059	0.085	Pass
Chamber size	−1.0019	0.4745	0.035	0.011	Fails on p
Appointed member	0.0100	0.0056	0.076	0.066	Pass

The failing cells are of two kinds. Gender and chamber size on the cross-branch pillar fail only the first condition: statistically significant at conventional levels while amount-

ing to imbalances of 0.087 and 0.011 standard deviations, the pattern large samples produce, in which precision rather than imbalance drives the p-value. Appointment status on the chamber pillar sits at the magnitude line (0.104 standard deviations, $p = 0.023$); section 7 notes that appointed and mid-term members never faced the close flip election the design exploits, and excluding them moves the premium by about one percent. Squire professionalism on the cross-branch pillar is statistically insignificant ($p = 0.20$) but carries the largest imbalance in the table at 0.164 standard deviations, and professionalism is one of the moderators in main-text Table 4. The state fixed effects absorb the between-state component of that imbalance and the professionalism-moderated estimates of section 5.2 engage what remains; we separate it out here because a magnitude failure and a p-value failure warrant different readings, and a pass count treats them alike.

The density test in the battery is run on the member-level running variable, which mirrors one chamber-level distance across all the members of a chamber. That test cannot detect the manipulation this design has to take seriously, which operates at the chamber level: a majority defending narrow control through candidate recruitment, targeted spending, or districting would sort chamber-terms, not members. We therefore run the test where the manipulation would occur. For each chamber-term with a known prior-term majority we sign the flip distance positive when the prior majority retained control and negative when control changed hands, which yields 940 usable chamber-terms across ninety-one chambers. If assignment near the threshold is as-if-random, retentions and losses must balance as the distance goes to zero, whatever the retention rate is away from it.

They balance within half a point and not beyond it. Within half a percentage point of the flip margin, fourteen chamber-terms are retentions and nine are losses (retention share 0.609; binomial $p = 0.41$ against a fair coin); within one point, thirty-four against nineteen (0.642; $p = 0.053$); within two points, sixty-nine against forty-five (0.605; $p = 0.031$). The corresponding log density ratios are 0.44, 0.58, and 0.43; the first interval con-

tains zero, the other two do not. Because a chamber that sits near the threshold across several terms contributes dependent chamber-terms, we also cluster by chamber ($p = 0.30, 0.021, 0.027$) and run a chamber-blocked sign-flip permutation that mirrors each chamber's whole pattern with probability one half ($p = 0.41, 0.051, 0.038$); the three methods agree, so the $\pm 1\%$ and $\pm 2\%$ bins fail the fair-coin test on all three. The mass is diffuse: at $\pm 2\%$ the 114 chamber-terms come from forty chambers, the five largest contribute 31 percent of the bin, and leaving any one chamber out moves the share only between 0.579 and 0.618. An era split places the rejection. Chamber-terms beginning before 2019, the window the earlier version of this analysis covered, pass at every bin (shares 0.579, 0.622, 0.583; $p \geq 0.12$ on all three methods), while the eighteen chamber-terms at $\pm 2\%$ from the 2019–2024 extension retain control at 0.722 (chamber-clustered $p = 0.008$ on twelve chambers; permutation $p = 0.058$; binomial $p = 0.096$). Two limits bound what the split can claim. The historical subset does not separate 0.58 from 0.5 for lack of power, so the point estimates in the two windows are not far apart; and a bin share cannot distinguish a continuous persistence slope from a discontinuity at zero, since any finite bin averages above one half under a positive slope. The direct test is the covariate discontinuity in lagged status reported in sections I and P: +0.336 at the chamber cutoff and +0.179 at the state cutoff. That imbalance exists, it accounts for the pre-LES placebo, and with lagged status as a control it moves the primaries by about 0.01. Whether the higher retention after 2018 is a sorting signal or the small-sample behavior of an eighteen-observation cell the data cannot say; a nationalization reading, in which narrow majorities have become more durable, is a hypothesis this test does not establish. The retention share by distance band, below, shows the shape: control is retained in essentially every chamber-term more than ten points from a flip, and the rate falls toward, without reaching, a coin flip at the threshold.

Distance to flip	Chamber-terms	Retentions	Share retained	Binomial p
Under 0.5%	23	14	0.609	0.41
0.5% to 1%	30	20	0.667	0.10
1% to 2%	61	35	0.574	0.31
2% to 5%	203	158	0.778	<0.001
5% to 10%	280	257	0.918	<0.001
Over 10%	343	340	0.991	<0.001

N Anzia-Jackman coding vintage: era-restricted gradient

The six Anzia-Jackman indicators arrive in the SLES bundle as a single cross-section, and the bundle does not record the year of the coding. No chamber in our panel takes more than one index value across the window — zero of ninety-one — so one coding is applied unchanged to chamber-terms three decades apart. Section 7 of the main text argues that the resulting misclassification attenuates the cartel-power gradient rather than manufacturing it. That argument is testable rather than merely assertable. If the coding describes late-window rules better than early-window ones, which is the direction any single-vintage coding of a changing institution runs, the gradient should be weaker in the early half of the window and no weaker in the late half.

Splitting the window at 2003 leaves 10,562 chamber cartel observations in 48 clusters before the split and 16,871 in 58 clusters after it. The continuous cartel-power slope is +0.132 in the early half ($p < 0.001$; wild bootstrap $p = 0.001$) and +0.185 in the late half ($p < 0.001$; wild bootstrap $p < 0.001$), against +0.161 estimated on the full window. The direction is the one the attenuation argument predicts, the gradient being some forty percent steeper in the half the coding is more likely to fit, but the two halves are not statistically distinguishable: a pooled model interacting treatment, the index, and a late-window indicator places the difference at +0.035 ($p = 0.11$; wild bootstrap $p = 0.17$). The cells are not

monotone within either half; the richest configuration is the highest in both.

Cartel-power index	1987–2002			2003–2024		
	Estimate	N	G	Estimate	N	G
0 (cartel-thin)	0.708	1,134	7	0.586	1,243	6
2	0.426	1,347	6	0.415	2,796	8
3	0.854	1,936	10	0.645	3,287	13
4	0.702	926	4	0.838	1,395	6
5 (cartel-thick)	0.961	5,219	21	1.268	8,150	25

Cells with fewer than twelve clusters use Webb six-point weights, following the convention of section H above; the cluster counts are reported so that the thinnest cells can be discounted accordingly.

The cross-branch pillar stays flat in both halves: +0.019 early ($p = 0.26$) and +0.041 late ($p = 0.14$), with an era difference of +0.037 ($p = 0.12$). The cross-branch pillar’s flat gradient holds within each half as well, so the proportional reading of section 5.2 is not an artifact of averaging a drifting coding over a long window. What the era split cannot do is date the coding, which the SLES bundle does not record; the caveat in section 7 accordingly stands as written, and the gradient should be read as an average over the window rather than as a rule-vintage-specific estimate.

O Stochastic forcing-variable rebuild

Both pillars’ running variables are deterministic uniform-shift constructions from observed vote shares; section 6.6 of the main text reports that rebuilding them in the stochastic style of Kirkland and Phillips (2020) — forty thousand simulated electoral-shock draws, recording for each member the smallest shock that flips the relevant condition (chamber-majority control for the chamber cartel pillar; party-specific unified government for the cross-branch pillar) — changes neither reading. This section reports

the estimates behind that statement. Each pair uses the pillar's primary fixed effects, controls and clustering at the five-percent bandwidth with tied observations excluded, the stochastic estimate obtained on the subsample where the simulated running variable is defined. The two constructions agree closely where both exist: the signed correlation between the deterministic and simulated running variables is +0.86 on the cross-branch pillar.

One feature of the stochastic cross-branch variable governs how its arm is estimated. Its sign follows the simulation's own assignment of unified status, which disagrees with the member's observed unified status for a fifth of the near-threshold sample (the two agree for 79.4 percent of member-terms within five percent). Treatment in the stochastic arm is the member's observed status, so the arm is a within-window contrast in observed status rather than an intercept shift at a cutoff the treatment follows, and the local-linear slope split does not apply to it in the same way. We therefore compare the two cross-branch constructions under the local-constant estimator, where both arms are within-window contrasts; for the chamber pillar, whose simulated variable's sign agrees with observed majority status, we report the comparison under both estimators.

Running variable	Estimator	Estimate	SE	N	G
<i>Chamber cartel pillar (μ)</i>					
Deterministic (primary)	Local-linear	0.818	0.108	27,433	67
Stochastic rebuild	Local-linear	0.713	0.093	14,142	57
Deterministic	Local-constant	0.802	0.063	27,433	67
Stochastic rebuild	Local-constant	0.760	0.083	14,142	57
<i>Cross-branch pillar (γ)</i>					
Deterministic	Local-constant	0.317	0.038	32,940	41
Stochastic rebuild	Local-constant	0.307	0.065	11,643	38

On the chamber pillar the stochastic estimate is 12.8 percent below the primary local-linear estimate and 5.3 percent below the local-constant one, with asymptotic intervals

that overlap in both comparisons (0.713 ± 0.181 against 0.818 ± 0.211); the stochastic arm passes thirteen of the fourteen cells of the bandwidth, fixed-effect, cluster, and donut grid of sections A–D, failing only the $\pm 1\%$ bandwidth (0.618 against 0.773). On the cross-branch pillar the two local-constant contrasts are 3.1 percent apart. Neither headline is an artifact of the uniform-shift assumption; the deterministic construction is retained for its larger analytical sample and tighter inference.

P Binding-institution decomposition by outcome, ratio-difference inference, and the lagged-status discontinuity

This section reports the three analyses that main-text sections 3, 4.1 and 4.2 draw on for the reading of the cross-branch pillar, under the primary local-linear estimator and, for disclosure, under the local-constant alternative. All three use the Pillar B primary specification (state and year fixed effects, state cluster, $\pm 5\%$ bandwidth, tied cases excluded) unless stated.

Part 1 decomposes each of the five outcomes by the binding institution at the cutoff, extending main-text Table 2 from the effectiveness score to the funnel components. Under the primary estimator every governor-binding cell, the design’s clean cross-branch estimate, is at or below zero and no wild interval excludes zero; the pooled bill_{law} estimate is carried by the own-chamber-binding cells, where both chamber-floor outcomes are also large and both substantive-significance outcomes are null, the chamber funnel. Under the local-constant estimator the governor-binding column is positive and two cells exclude zero, bill_{law} among them; that cell is the one reading in the paper that the choice of estimator decides, and section 3 gives the grounds for the primary choice.

Binding institution	Outcome	Local-linear (primary)		Local-constant	
		Estimate (SE)	Wild 95% CI	Estimate (SE)	Wild 95% CI
Governor (N 18,939; G 40)	LES	-0.075 (0.060)	[-0.218, 0.066]	0.141 (0.032)	[0.064, 0.220]
	bill_pass	-1.129 (0.611)	[-2.785, 0.450]	0.462 (0.584)	[-0.805, 1.835]
	bill_law	-0.068 (0.416)	[-1.173, 1.060]	0.706 (0.279)	[0.060, 1.377]
	ss_pass	-0.123 (0.148)	[-0.458, 0.219]	0.098 (0.080)	[-0.083, 0.270]
	ss_law	0.062 (0.116)	[-0.198, 0.309]	0.133 (0.052)	[0.024, 0.240]
Own chamber (N 6,650; G 31)	LES	0.672 (0.117)	[0.357, 1.023]	0.744 (0.072)	[0.541, 0.963]
	bill_pass	3.860 (1.488)	[0.060, 7.431]	4.401 (1.098)	[1.904, 7.056]
	bill_law	3.421 (1.133)	[0.657, 6.046]	3.120 (0.710)	[1.489, 4.904]
	ss_pass	0.104 (0.393)	[-0.870, 1.060]	0.303 (0.126)	[0.031, 0.565]
	ss_law	0.146 (0.315)	[-0.642, 0.918]	0.269 (0.101)	[0.053, 0.479]
Other chamber (N 7,351; G 31)	LES	-0.051 (0.111)	[-0.292, 0.210]	0.264 (0.078)	[0.051, 0.472]
	bill_pass	0.826 (1.108)	[-2.090, 3.816]	1.898 (0.795)	[-0.379, 3.961]
	bill_law	1.427 (0.762)	[-0.453, 3.315]	2.235 (0.557)	[0.633, 3.874]
	ss_pass	-0.048 (0.133)	[-0.361, 0.251]	0.135 (0.077)	[-0.037, 0.310]
	ss_law	0.096 (0.118)	[-0.207, 0.392]	0.151 (0.051)	[0.034, 0.271]

Part 2 attaches inference to main-text Table 3. The two pillars' analytical samples are resampled jointly by state (forty-seven states in the union of the two samples, thirty-two common to both; 999 replications), both pillars re-estimated on each draw, and each ratio difference recomputed; the joint Wald test asks whether the four differences are zero, which is what a cross-branch pillar that carried only the chamber premium would produce. Under the primary estimator no difference excludes zero and the joint test does not reject ($\chi^2(4) = 1.46$, $p = 0.83$); under the local-constant estimator the bill_law difference excludes zero and the joint test rejects ($\chi^2(4) = 19.2$, $p < 0.001$). The cross-branch effectiveness estimate is the denominator of every cross-branch ratio, so under the primary estimator (0.198, standard error 0.059) the bootstrap ratios are unstable and the intervals wide; the non-rejection is consistent with equal profiles and with much else, and is not evidence of equality.

Outcome	Local-linear (primary)		Local-constant	
	Ratio diff	95% interval	Ratio diff	95% interval
bill_pass	-2.00	[-7.69, 2.62]	-0.68	[-3.21, 1.26]
bill_low	+2.41	[-0.17, 10.43]	+1.60	[0.17, 3.01]
ss_pass	-0.07	[-1.62, 1.53]	-0.07	[-0.41, 0.24]
ss_low	+0.45	[-0.58, 2.26]	+0.19	[-0.03, 0.48]
Joint Wald $\chi^2(4)$, p	1.46, p = 0.83		19.2, p < 0.001	

Part 3 is the covariate discontinuity in lagged status that sections I and M rely on: the lagged-majority indicator (chamber pillar) and the lagged-unified indicator (cross-branch pillar) regressed on current treatment at each pillar's cutoff, on the samples with a defined lag (19,962 and 24,084 member-terms). The product of each primary premium and the measured jump is the placebo that persistence alone predicts; the measured pre-LES placebo is shown beside it.

Pillar	Estimator	Jump (SE)	Wild 95% CI	Premium $\times$ jump	Placebo
Chamber	Local-linear	0.336 (0.131)	[0.028, 0.634]	0.275	0.279
Chamber	Local-constant	0.431 (0.053)	[0.313, 0.548]	0.346	0.369
Cross-branch	Local-linear	0.179 (0.107)	[-0.045, 0.403]	0.035	0.120
Cross-branch	Local-constant	0.449 (0.057)	[0.302, 0.583]	0.143	0.172

Under both estimators the chamber-pillar arithmetic closes to within 0.02, and under the local-constant estimator the cross-branch arithmetic closes to within 0.03; under the primary estimator the cross-branch placebo (0.120, standard error 0.084) is itself not distinguishable from zero, so the looser fit there is inside its own noise. The lagged-status imbalance that the chamber-term bin shares of section M imply (0.28 at $\pm 1\%$, 0.21 at $\pm 2\%$) is smaller than the jump measured here; section I reports that controlling for lagged status moves the primaries by about 0.01.