

Prices Without Rents: What an Election Delivers (and Doesn't) to Donor Firms

Kisoo Kim*

August 26, 2026

Abstract

Do corporate campaign contributions purchase measurable government benefits for donor firms? Using a regression-discontinuity design on 1,609 publicly traded firms linked through corporate PAC contributions to 588 close U.S. House general elections (2000–2022), we test four government-benefit outcomes and three political-behavior outcomes, each against a matched pre-trend placebo. No channel delivers a benefit to the donor firm. Three of the four government-benefit outcomes – procurement allocations, effective tax rates, and SEC enforcement exposure – return null estimates against clean pre-trend placebos; the fourth, federal grants, returns a small positive estimate whose own placebo is borderline. Two of the three political-behavior outcomes are likewise null, and the third is small, statistically significant, and signed opposite to the patronage prediction: winner-supporting firms lobby slightly less, not more, on the newly connected legislator's committee jurisdictions. The stock-market reaction to a donor-favored win, an order of magnitude smaller than prior estimates in the literature, fails a standard bandwidth-and-polynomial robustness grid, and even under its most generous specification does not persist beyond short horizons. Corporate bond spreads show no compression for donor-connected winners either. Our paper joins a growing null-findings literature on corporate campaign-contribution returns.

Keywords: political connections, close elections, regression discontinuity, campaign finance, corporate political strategy, event study

*Email: kisoo@virginia.edu

1 Introduction

Corporate spending on U.S. federal political campaigns runs into the billions of dollars each election cycle and persists across decades of corporate life. The scale and persistence of that spending pose a question: what do firms receive in exchange? A substantial body of research finds that political connections move share prices when those connections form, dissolve, or are revealed. Close elections in which a firm’s favored candidate narrowly wins generate positive cumulative abnormal returns on the firm’s stock (Akey, 2015); politically connected firms command measurable value premia around connection-changing events (Faccio, 2006; Goldman et al., 2009); and aggregate legislator voting behavior predicts firm-level stock returns in the direction that favors firms with identifiable policy stakes (Cohen et al., 2013). This “connections-matter” literature is often read – explicitly or implicitly – as evidence that connections deliver tangible benefit: procurement contracts, favorable tax treatment, reduced enforcement, credit-market advantage, or informational access that the market subsequently prices. Yet no empirical outcome test has directly established that link.

This paper conducts the missing outcome-side test. Extending the close-election regression-discontinuity design of Akey (2015) to twelve U.S. House general-election cycles (2000–2022; the 2024 cycle is covered by our panel but excluded from causal regressions because its post-window is truncated at the time of analysis), we assemble a panel of 1,609 publicly traded firms linked through corporate PAC contributions to House candidates, with 588 close general elections – those decided by a margin of victory within five percentage points – as the causal sample. On this panel we run ten analyses organized into three axes. We first test four direct government-sourced benefit outcomes: federal procurement allocations, effective tax rates (ETR), SEC enforcement actions captured in Accounting and Auditing Enforcement Releases (AAER), and federal grants. We then test three political-behavior outcomes, a revealed-preference consistency check on how firms adjust their own post-election political investment: Lobbying Disclosure Act

expenditure, lobbying target alignment with the winning legislator’s committee jurisdictions, and PAC contribution strategy measured as cycle-over-cycle portfolio changes. Finally, we run three market-side analyses: we replicate Akey (2015)’s event-window cumulative abnormal return under the standard bandwidth $\times$ polynomial robustness grid, extend it to a long-horizon abnormal-return analysis out to one trading year, and test corporate bond spreads as a credit-market analog to the equity reaction.

The evidence is coherent across outcome channels. On the government-benefit axis, three of four outcomes are null – procurement, effective tax rates, and SEC enforcement actions – with confidence intervals that rule out the effect magnitudes the adjacent literature documents. The fourth, federal grants (grants-only, excluding loans and direct payments), returns a small positive estimate ($\approx 5\%$ log-scale, $p = 0.041$) whose own pre-trend placebo is borderline ($p = 0.13$); a single coefficient at $p = 0.041$ in a nine-outcome family is within the false-positive count that many tests produce by chance, and no conclusion in the paper rests on it.

On the political-behavior axis, all three outcomes are null or signed opposite to the patronage prediction: lobbying expenditure is a precise null, with a point estimate about nine percent below the loser-side baseline on the log scale ($p = 0.14$); lobbying on the winner’s committee-jurisdiction issue codes is small but statistically negative (approximately -4% log-scale, $p < 0.001$, though the corresponding pre-period placebo is marginal at $p = 0.09$); and PAC contributions show portfolio expansion rather than insider concentration: firms that donated to the winning candidate before the election are, in the next cycle, 23 percentage points more likely to donate to him again than firms that had donated to his opponent – down from a 29-percentage-point gap before the election, rather than widening as a patronage reading would predict.

On the market-reaction axis, a direction-consistent point estimate is recovered in a subset of specifications, but the magnitude is an order of magnitude smaller than Akey’s and the estimate does not survive the bandwidth $\times$ polynomial grid: 20 of 32 grid cells

are negative and only one specification reaches conventional significance. Even under the most generous specification choice, the initial reaction does not persist: at three of five post-event horizons the 95% CI upper bound falls below the initial magnitude. The RDD on corporate bond spreads is null, with a point estimate of about 2.4 basis points of compression ($p = 0.25$) and a 95% confidence interval that rules out compression larger than roughly seven basis points.

A single reading organizes these patterns. Whether markets price a connection premium at election revelation is empirically contested; regardless, within our horizon and on the outcomes we measure, realized fundamentals and firms' own political investment do not differ between donor-connected winners and comparably close losers. If there is a premium, it is not subsequently validated by delivery; if there is no premium, the absence of realized benefit is unsurprising but remains a non-trivial finding against the connections-matter literature's implicit claim. Firms, for their part, do not behave as if the connection has given them new leverage; if anything, their post-win political investment diversifies. Perhaps the market-reaction literature has been documenting information updating about legislator identity and committee composition, rather than anticipation of realized rents that subsequently flow to donor firms.

The paper contributes to two connected conversations. First, it challenges the close-election CAR literature on two fronts. The event-window reaction itself is not robust on our larger sample, reinforcing Fowler et al. (2020)'s critique of the original finding; and, independent of whether any initial reaction is accepted, the realized fundamentals on which such an expectation would have to be grounded do not differ between donor-connected winners and comparably close losers. We provide the missing outcome-side evidence and find it largely null, complementing Fowler et al. (2020)'s close-election RDD on firm market value and Fourinaies and Fowler (2021)'s difference-in-differences on state campaign finance regulations; across these distinct designs, the three papers together suggest that prior positive findings may have been driven by selection or composition rather

than causal treatment. Second, we substitute a portfolio account for a patronage account of corporate political strategy. Firms behave strategically in ways the literature documents, yet in the close-election slice where that strategic investment most plausibly pays off, they neither concentrate giving on the winner nor lobby more on his committee’s jurisdiction – behavioral evidence consistent with portfolio diversification rather than failed patronage.

2 Instrumental Exchange and the Market-Pricing Step

This paper takes as a working premise the view – well represented in the political-economy literature on corporate political activity – that firms donate to political campaigns as instrumental investment rather than consumption. Shareholder-financed PAC contributions to partisan candidates are difficult to justify under a pure ideological or expressive rationale, and the empirical regularities of corporate giving – heavy concentration in regulated industries (Ansolabehere et al., 2003), strategic allocation toward committee members with jurisdictional authority over the donating firm (Grier et al., 1994; Snyder Jr., 1992), and cyclical adjustment to anticipated electoral outcomes (Bonica, 2016) – line up much better with an exchange motive than with a consumption motive. Taking this investment premise as given, we ask whether the market’s pricing at the moment a close-election win is called is subsequently validated by realized firm fundamentals and by firms’ own political behavior. The unit of expectation in our design is the market; the unit of validation is the firm.

A positive cumulative abnormal return is, in principle, a forecast that integrates across three channels through which the connection might subsequently transfer value. The first is a direct transfer from the government to the firm: procurement contracts allocated on favorable terms (Goldman et al., 2013; Agca and Igan, 2023; Baltrunaite, 2020), effective tax burden reduced through statute or audit, crisis-era government investment alloca-

tions (as documented for the TARP Capital Purchase Program by Duchin and Sosyura 2012), or enforcement leniency from regulators whose oversight committees are staffed by connected legislators (Correia, 2014). The second is an indirect, market-mediated benefit: political access may reduce the firm's perceived risk of adverse regulation and therefore compress its cost of capital, a pattern documented for firm valuations (Faccio, 2006) and for firms' financing decisions (Boubakri et al., 2012). The third is informational, in the sense of access to the legislative calendar and private signals about the direction of policy; this channel is harder to observe directly but has received prominent attention in the political-connections literature that links legislator-level information to firm-level cash-flow outcomes (Cohen et al., 2013), with the more contested question of whether legislators themselves earn abnormal returns on their own portfolios an active dispute – claimed by Ziobrowski et al. (2004) and rebutted by Eggers and Hainmueller (2013). Positive cumulative abnormal returns on donor-firm stock around events that resolve political uncertainty – documented by Akey (2015) in a close-election RD setting and in related event-study designs going back at least to Jayachandran (2006) – are, in this framing, a statement about what the market anticipates the connection will be worth, integrated across all three channels. The sign tells us the market thinks the connection is valuable. The magnitude tells us how much.

What a positive market reaction does not directly tell us is whether the anticipated benefit subsequently materializes. An efficient market prices what it expects to happen; whether what actually happens confirms, exceeds, or falls short of the priced expectation is a distinct empirical question. If the market's pricing at the moment a close election is called correctly forecasts realized connection-related benefit, we should observe measurable changes in firm-level fundamentals: on the government side, more federal contracting, lower effective tax rates, fewer SEC enforcement actions, or more grant dollars; on the credit-market side, compression in corporate bond spreads reflecting reduced perceived risk – concentrated on firms whose donated candidate narrowly won relative to

firms whose donated candidate narrowly lost. The close-election comparison isolates the marginal effect of the connection from the pooled effect of the donation: both sets of firms donated; only one set saw their candidate win.

A second, complementary prediction concerns firms' own subsequent political behavior. If firms believe a close-election win has increased their political leverage, they should respond by investing more in exploiting that leverage. Three observable behaviors are natural candidates. Lobbying expenditure should rise, because the marginal dollar of lobbying is worth more when a sympathetic legislator is in place to receive it. The target composition of lobbying should shift toward the policy domains and issue codes associated with the newly connected legislator's committee jurisdictions, because targeted lobbying returns exceed untargeted lobbying returns in a principal-agent framework where the connected legislator is the relevant agent (Hall and Wayman, 1990; Bertrand et al., 2014). PAC strategy should shift toward concentration on the newly connected legislator, because subsequent donations operate as maintenance payments on an established relationship. If, by contrast, firms observe that the connection does not translate into measurable advantage on the channels they can monitor, they have no instrumental reason to double down. Their behavior then ought to look much like the behavior of firms whose candidate narrowly lost. This is the revealed-preference logic we rely on.

The two predictions generate a four-cell taxonomy. If the market reaction moves and government-benefit outcomes move along with it, the reading that connections deliver is supported. If the market reaction moves but government-benefit outcomes do not, yet firms increase their political investment, a weaker version holds: firms believe in the connection even though realized delivery lies below the measurement floor of our outcome data. If the market reaction moves and neither government-benefit outcomes nor firms' political investment move, the priced expectation is best read as information pricing: the market prices an expectation that does not materialize and that firms themselves, given more time and better information than the market had at the moment of revelation, do

not act on as if it were a live asset. Finally, if the event-window market reaction itself is not robust to standard specification variation, both the information-pricing and the realized-delivery readings lose their premise: the connections-matter empirical claim on the close-election contrast then fails on two margins at once, the expectation-forming step and the realization step alike, and the absence of realized-outcome movement is then itself the substantive finding, with no priced-expectation frame needed to give it content. As §5 and §6 will show, this fourth cell is where our panel's evidence lies.

Patronage and Portfolio Accounts of Political Investment

The revealed-preference predictions above – lobbying expenditure up, lobbying composition shifted toward the connected legislator's jurisdictions, PAC giving concentrated on the winner – follow from a specific account of what a political tie is. Call it the patronage account: a tie is a relationship-specific asset whose value attaches to the identity of one legislator, is built through repeated exchange, and is maintained by continued investment in that legislator (Hall and Wayman, 1990; Bertrand et al., 2014). On this account a close-election win raises the return to every dollar directed at the newly seated legislator, so the optimal post-win response is concentration on the winning relationship. A competing account, which we call the portfolio account, treats ties as partially substitutable instruments for access: the firm holds positions across many legislators and both parties (Fournaies and Hall, 2018; Barber, 2016), donations and lobbying operate as linked components of a single influence strategy (Kim et al., 2026), and one legislator's electoral fate changes the value of the position set only marginally. On this account a close-election win resolves one position, and the optimal response is rebalancing: total political expenditure approximately flat, recipient breadth maintained or extended, loyalty to the specific winner not amplified, and, for a firm hedging against future majorities, giving extended toward the party that just lost the seat. The two accounts agree that corporate giving is strategic; they separate on what a marginal win does to subsequent behavior. The three political-

behavior outcomes tested in §7 – lobbying expenditure, lobbying target composition, and PAC allocation – measure exactly the margins on which they separate.

The Strongest Case for Delivery

The setting we study is the one in which realized-delivery effects should be largest. Three conditions jointly make this the most-likely case for delivery. First, we restrict to publicly traded Compustat firms that contributed through a corporate PAC to a specific close-election candidate in the pre-cycle window – the subpopulation the existing event-window literature (Akey, 2015; Goldman et al., 2009; Cohen et al., 2013) identifies as both the most measurably connected and the one in which the largest putative connection-value effects have been documented. Second, close-election resolution removes ex-ante uncertainty about whether the connected candidate wins, isolating the realized-win treatment from the pooled donation decision. Third, our design is powered to detect the realized-benefit magnitudes the adjacent literature has documented: minimum detectable effects on the RDD specification lie below the canonical effect sizes in Akey (2015), Agca and Igan (2023), Baltrunaite (2020), and Correia (2014) (see §6.4). As Fourniaies and Fowler (2021) argue in the structurally parallel insurance-industry setting, a null on realized delivery obtained in a most-likely-case sample is informative about the general population of firm-legislator connections rather than about a corner case. If realized delivery is present anywhere, it should be present here.

Two boundaries on the theoretical frame deserve brief acknowledgment. First, the exchange framework does not require that every donating firm be strictly quid-pro-quo; it requires only that donation decisions be made with some expectation of return and that those expectations aggregate across the donor universe. Heterogeneity in firm motives is compatible with the null findings we will report. Second, the validation horizon matters. Our analyses measure outcomes on short-to-medium-run post-election horizons: two fiscal years after the election for most outcomes, and two election cycles (roughly four years)

for PAC strategy. Effects that operate on longer horizons – goodwill built over multiple election cycles, quiet regulatory posture shifts that take years to manifest in observable enforcement data – are not within our detection range. We take this as a real limitation and return to it in §9. The CAR literature does not specify when priced benefits should materialize – a three-day price move can discount value flows that take years to arrive – so an asymmetry between a short pricing window and a longer validation horizon is a real possibility we cannot rule out.

3 Data

3.1 Panel Construction

The analysis panel is a firm-candidate-election panel covering all U.S. House of Representatives general elections from 2000 to 2024, a span of thirteen election cycles. Corporate PAC contributions from firms linked to Compustat GVKEYs are obtained from the FEC’s `pas2*` bulk files; candidates and committees are identified through the FEC `cn*` and `ccl*` files. PAC-to-GVKEY linkage comes from Dane Christensen’s `PAC_GVKEY_Linktable_1991_2020` (Christensen et al., 2022, 2023), which we extend through 2024 (Appendix A). House general-election outcomes and margin-of-victory calculations come from the MIT Election Lab’s 1976–2024 returns file. The assembled panel contains 215,332 firm-candidate-election rows covering 1,609 distinct firms. The 2024 cycle carries a truncated post-window and is excluded from the main pooled regressions.

3.2 Main Analysis Sample

The main sample restricts to close elections with a two-party margin of victory at or below five percentage points, the bandwidth used by Akey (2015) and several adjacent studies. The close-5pp restriction yields approximately 588 elections and approximately 35,000

firm-election observations on the 2000–2022 cycles. Following Akey, we further exclude firms that donated comparable amounts to both candidates in a given close race: for these firms no causal contrast exists at the electoral margin. The close-race hedging rate has declined sharply from 4.3% in the 2010 cycle to 0.6% in the 2022 cycle, so this restriction drops a small and shrinking share of observations. As a robustness check we also re-estimate every primary outcome on the full close-race sample, retaining firms that donated to both candidates and assigning them to the winner or loser side by the sign of their net donation (Appendix F.7); the direction and significance of every outcome are preserved.

3.3 Outcome Data

Counts reported in this subsection are panel coverage: the firm-years each source supplies for the firms in our panel. The estimation sample for a given regression is a count of close-5pp firm-elections and is reported with the estimate itself.

Cumulative abnormal returns. Daily Center for Research in Security Prices (CRSP) returns combined with Fama–French three-factor daily factors produce firm-level abnormal returns on a $(-1, +3)$ trading-day window around the close-election date, using a $[-130, -11]$ estimation window. We denote the three-day cumulative reaction $CAR_3 = CAR[-1, +1]$. The CAR panel covers generals 2000–2020 (11 cycles); it pre-dates the 2022/2024 panel extension and is not refreshed here because the substantive CAR finding is demoted in this draft (§5). For the persistence analysis we construct eight post-event windows extending to 252 trading days and three pre-event placebo windows.

Federal procurement. Prime-contractor obligations from USASpending’s Federal Procurement Data System (FPDS) are aggregated by GVKEY-fiscal-year. The panel covers

FY 2000 through FY 2026 (partial); firm-years with no observed obligation are carried as zeros.

Effective tax rate. GAAP effective tax rate is computed as tax expense over pre-tax income from Compustat, winsorized to $[0, 1]$. The panel contains 29,507 firm-years on 1,573 firms.

Bond spreads. Yield-to-maturity spreads to comparable-duration Treasury come from WRDS, aggregated to the firm-year. The issuance-amount-weighted spread measure covers 12,945 firm-years on 969 firms (61% of the main sample; the public-debt-issuer subsample), winsorized to $[0, 0.05]$.

SEC enforcement. Accounting and Auditing Enforcement Releases (AAER) come from WRDS. The panel identifies 155 ever-flagged firms and 706 violation-period cells across 42,579 firm-years. The primary outcome is an indicator for any violation event within a measurement window.

Federal grants and assistance. Non-procurement outlays from USASpending's Financial Assistance Broadcast Service (FABS) cover 674 of 1,609 main-sample firms over FY 2000–2024 and total approximately \$69 billion. We report three variants: total assistance, any-assistance indicator, and grants-only (excluding loans, loan guarantees, direct payments, and insurance). The grants-only variant is the cleanest substantive measure of direct budget allocation.

Federal lobbying. Lobbying Disclosure Act filings come from LobbyView, aggregated from quarterly client disclosures over 1999–2024. The panel contains 29,543 firm-years on 3,074 ever-lobbying firms. Variants include log-dollar expenditure, the any-lobbying indicator, and the log number of reports.

Committee-issue-code crosswalk. Winner committee assignments come from Charles Stewart’s House Committee Assignments file, recovered for pre-2022 cycles via a join to the Voteview HSall member file. A hand-coded crosswalk maps each of the eighteen standing House committees to a set of LDA issue codes (Armed Services $\rightarrow$ {DEF, AER, HOM}; Financial Services $\rightarrow$ {FIN, BNK, INS}; and so on). The crosswalk underlies the lobbying-target-alignment outcome in §7.

3.4 Outcome Windows

For outcomes observed at fiscal-year granularity we define the pre window as fiscal years $\{t-1, t\}$, the post window as $\{t+1, t+2\}$, and the earlier-pre placebo window as $\{t-3, t-2\}$. The post window $\{t+1, t+2\}$ corresponds to the winner’s full two-year House term (January of year $t+1$ through January of year $t+3$); for firms with calendar fiscal years this coincides exactly with the term, and for federal fiscal-year outcomes (procurement) the window contains one to two months of pre-inauguration activity, which attenuates rather than inflates estimated effects. For PAC contributions, which cycle every two years, we use cycle-based windows: pre is cycle $t-2$, post is cycle $t+2$, placebo is cycle $t-4$.

4 Identification

Our design improves on the existing literature in three respects that together yield a sharper test of the close-election connection-value hypothesis than prior work has delivered. First, prior close-election studies of donor-firm procurement rely on level regressions that are vulnerable to between-firm composition when the outcome is size-correlated. On our own panel a pre-election level placebo is statistically indistinguishable from the post-election level coefficient, identifying the level estimate as composition rather than causation. We use instead a within-firm Δ -RDD that absorbs composition under a testable pre-trend placebo. Second, where existing work typically reports a sin-

gle inferential frame, we triangulate across clustered OLS, local-polynomial `rdrobust` (Calonico et al., 2014), and Fisherian local-randomization `rdrandinf` (Cattaneo et al., 2016, 2020); disagreements across frames are themselves diagnostic. Third, where existing work pools without a size-correlation check, we deploy the pre-election level placebo as a cheap, general-purpose diagnostic that any close-election RDD on firm-level outcomes can run on its own panel. These three improvements matter precisely because prior close-election findings on donor-firm procurement are (as we show below) compatible with composition bias rather than causal delivery, so a cleaner test of the same claim is informative even when it returns a null.

4.1 Close-Election RDD and the Level-Regression Pitfall

The close-election regression-discontinuity design, formalized by Lee (2008) and used in Akey (2015), Agca and Igan (2023), and Baltrunaite (2020), with related discontinuity-design evidence on agency sign-off thresholds in Fazekas et al. (2022), exploits the fact that near a 50–50 race the identity of the winner is plausibly as-if random conditional on a narrow electoral margin. Under the standard continuity assumption, firms whose candidate narrowly wins are comparable in expectation to firms whose candidate narrowly loses.

Recent methodological work has converged on what close-election designs identify and when their estimates deserve trust, and the composition problem we document is distinct from the failure modes that literature describes. Marshall (2024) shows that when the treatment is a winning politician’s characteristic, the design bundles the characteristic of interest with every correlated characteristic, including the compensating differentials that let such candidates reach close races at all; Bertoli and Hazlett (2025) clarify that such designs identify a district-level effect of electing one candidate type rather than the isolated effect of the characteristic; and De Magalhães et al. (2025) benchmark close-election RDD estimators against elections resolved by actual lotteries, establishing which specifi-

cations recover experimental answers. All three critiques operate on the candidate side. In our design the treatment is not a candidate characteristic: both sides of the cutoff elect a winner, and the contrast is between firms that funded the winner and firms that funded the loser. The selection that threatens this design is therefore on the donor side (which firms chose to fund which races), and because many firms stack onto each race, imbalance concentrates in firm-level covariates such as baseline contractor size rather than in district-level ones. The candidate-side critiques bound what any close-election contrast can identify; the pre-period level placebo we develop below diagnoses a failure mode specific to firm-level designs on size-correlated outcomes, and the two sets of diagnostics are complementary rather than substitutable.

A level regression of post-window procurement on the win indicator at $|\text{margin}| \leq 0.05$ returns $\hat{\beta} = +0.308$ on the log scale, with cluster-robust $p = 0.013$. The identical regression with the *pre-election* procurement window as the dependent variable returns $+0.301$ with $p = 0.022$. The two coefficients are statistically and substantively indistinguishable: whatever gap exists between winner-supporting and loser-supporting firms on post-election procurement was already present before the election. The design has selected a sample in which winner-donor firms are systematically larger contractors at baseline, and the level-RDD coefficient picks up that standing difference rather than any causal effect of the election. A paired-permutation balance test confirms the diagnosis: 98.5% of permutation draws return positive pre-period imbalance on the pooled sample. Fisherian `rdrandinf` on the same level specification returns $p = 0.82$ – more than an order of magnitude larger than the cluster-robust p -value reported above. Disagreement across inferential frames of that magnitude is itself a signal that the design, as applied in level regressions, is not identifying a treatment effect.

4.2 The Δ -RDD and the Pre-Trend Placebo

If the binding problem is a firm-specific level difference between winner-supporting and loser-supporting donors, differencing within firm should absorb it. Define the within-firm-election change:

$$\Delta y_i \equiv \bar{y}_{i,\text{post}} - \bar{y}_{i,\text{pre}}, \quad (1)$$

where the windows are as specified in §3 and y is an outcome-specific transformation of the raw measure. The identifying regression is

$$\Delta y_i = \alpha + \beta \cdot \text{win}_i + \varepsilon_i, \quad (2)$$

estimated with election-clustered standard errors at $|\text{margin}| \leq 0.05$, excluding hedgers.

Differencing absorbs level composition but cannot absorb pre-existing trends. To detect trend contamination we run an analogous regression on the earlier-pre window:

$$\Delta y_i^{\text{placebo}} \equiv \bar{y}_{i,\text{pre}} - \bar{y}_{i,\text{prelag}}, \quad \text{prelag} = \{t-3, t-2\}. \quad (3)$$

Placebo nullity is a necessary condition for causal interpretation of any Δ -RDD coefficient we report.

Applied to procurement, the pooled 2000–2022 Δ -RDD returns $\hat{\beta} = +0.006$ with standard error 0.045 and $p = 0.88$; the placebo returns +0.073 with $p = 0.26$. Levels are composition-contaminated; the pre-trend placebo is clean; the differenced coefficient is therefore defensibly causal and small. A 95% confidence interval of approximately $[-0.08, +0.09]$ on the log scale rules out post-election procurement effects larger than about ten percent. This is a tight bound: it rules out the magnitude of Akey’s (2015) level-RDD and of the adjacent procurement-connections literature’s effect sizes on the candidate-level electoral margin.

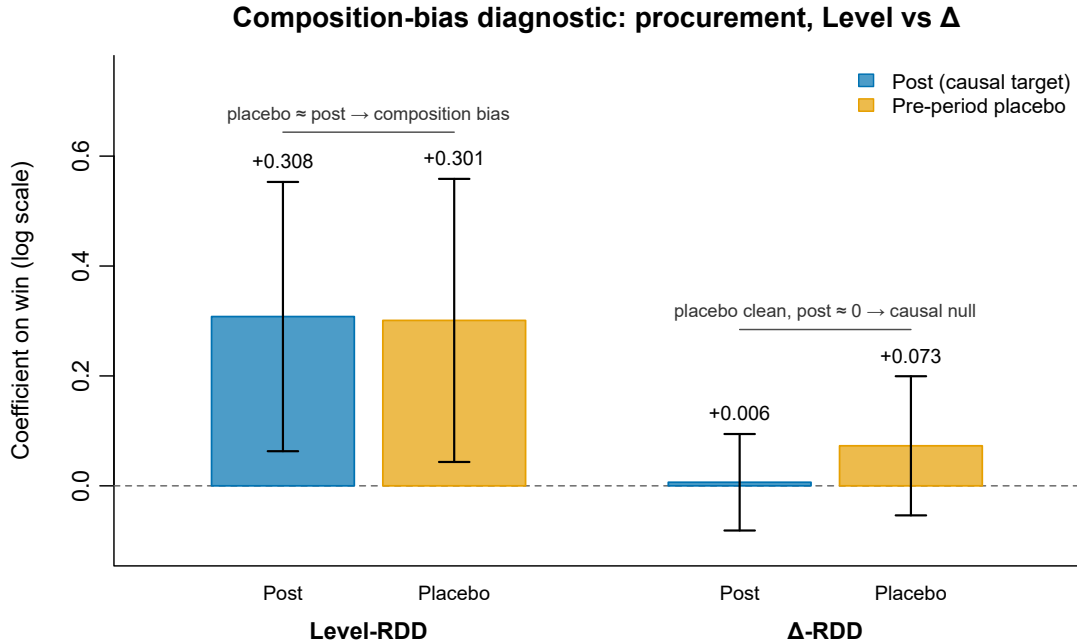


Figure 1: Composition-bias diagnostic on procurement. The level-RDD returns +0.308 on the post window and a statistically and substantively identical +0.301 on a pre-election placebo window, showing that the level estimate tracks a standing between-firm size difference rather than any causal effect of winning. The within-firm Δ -RDD returns a tightly estimated +0.006 on the post window against a clean placebo of +0.073, indicating no causal delivery. Bars are coefficients on the win indicator at $|\text{margin}| \leq 0.05$; whiskers are 95% cluster-robust confidence intervals.

4.3 Estimator Triangulation

The running variable in a close-election RDD is the margin of victory, which in our panel exhibits heavy repeated values at many discrete margins: roughly 555 distinct margin values across 33,000 firm-election rows. The continuity framework underlying the default `rdrobust` local-polynomial estimator assumes a continuous density at the cutoff, which discrete margins violate mechanically (Cattaneo et al., 2024). For rare binary outcomes the pathology is more acute: `rdrobust` on the AAER violation indicator in the post-CU subsample returns a t -statistic of -73.8 with a standard error of 0.0002 , a classic degenerate estimate; the rarer AAER release indicator is degenerate in the same subsample in the opposite direction, at $t = 9.3$. Following Cattaneo et al. (2024), we report all three estimators – clustered OLS, `rdrobust`, and Fisherian `rdrandinf` – on the main

Δ -RDD specification and on the placebo, for each outcome and each subperiod (pooled 2000–2022, pre-CU 2000–2008, post-CU 2010–2022). Agreement across frames is robust evidence; disagreement is diagnostic.

A related question is whether an MSE-optimal bandwidth following Calonico et al. (2014)’s CCT-MSERD procedure would be preferable to the fixed 5pp window we use throughout. Running CCT-MSERD on procurement yields $h_{\text{opt}} \in [0.023, 0.049]$ across subperiods – always narrower than the 5pp window. We retain 5pp as primary for three reasons. First, the narrower h_{opt} reduces the effective sample without moving the substantive estimate: at $h_{\text{opt}} = 0.036$ pooled OLS returns $\hat{\beta} = +0.006$ with $p = 0.91$, essentially identical to the 5pp estimate. Second, adopting h_{opt} as primary elevates `rdrobust` – the frame that selected it – to the dominant inferential role, which conflicts with the mass-point pathology documented in the preceding paragraph. Third, `rdwinselect` returns no balanced window at any width on this panel, so the Fisherian leg of our triangulation cannot operate at h_{opt} either. Preserving 5pp keeps all three estimators live, satisfies the triangulation rule, and leaves the substantive null unchanged.

4.4 Cross-Race Aggregation

A further design issue arises because treatment is assigned at the race level while outcomes are measured at the firm level. A firm that funds several close races in the same cycle carries a mixture of wins and losses into a single Δy_i : in our panel the median close-race firm-cycle touches three close races, and 55.9% of close-race firm-cycles contain both winning and losing bets. Pooling across races therefore attenuates any per-race treatment effect toward zero mechanically, whether or not connections deliver. Three responses bound this concern. First, we re-estimate the Δ -RDD on the firm-cycles that contain exactly one close race, where no mixture is possible by construction. Because a firm may fund many races of which only one is close, this slice retains 1,886 firm-elections from 986 firms and 468 elections, roughly nine times the single-candidate slice of Appendix F.5.

Procurement on the slice is -0.174 ($p = 0.19$) with a clean placebo ($p = 0.88$) and a 95% confidence interval of $[-0.44, +0.09]$: less precise than the pooled estimate, but negative in sign and still excluding the $+0.28$ -order procurement effects the adjacent literature documents. None of the four government-benefit outcomes shows delivery on the slice; the full nine-outcome results, including the two cells that require qualification, are in Appendix F.6. Second, if aggregation attenuation drove the pooled null, estimated effects should strengthen as portfolio concentration increases; the breadth-stratified estimates in Appendix F.5 show no such gradient on any outcome. Third, the within-firm fixed-effects specification (§8), which estimates the win contrast within firm across races and years and is therefore least exposed to donor-side composition, returns a pooled procurement coefficient of -0.002 with a 95% confidence interval of $[-0.05, +0.05]$, converging on the same tight zero from a differently structured comparison.

4.5 Window Convention and Sample Restrictions

Three restrictions run throughout: $|\text{margin}| \leq 0.05$ for the close-election sample; exclusion of hedging firms; and exclusion of cycles with incomplete post-windows (principally 2024). The pre-CU versus post-CU split uses 2000–2008 as pre-CU and 2010–2022 as post-CU, placing the January 2010 *Citizens United v. FEC* decision at the cutoff; the two subperiods partition the pooled 2000–2022 sample exactly. The cumulative-abnormal-return panel is the one exception: it ends with the 2020 cycle (§3), so its post-CU subperiod is 2010–2020. Robustness to alternative restrictions – tighter bandwidths, retaining hedging firms by sign of net donation, and within-firm fixed effects on the dual-appearance subpopulation (Appendix F) – is reported in §8 and at greater length in the appendix. The paired-permutation balance test and the four-cell level-versus- Δ decomposition are reported in Appendix B, and the special-election augmentation in Appendix F; together they form a compact transportable diagnostic for firm-level close-election RDD on size-correlated outcomes.

5 Market Reaction Does Not Replicate

5.1 Event-Window Estimate under the FGS Robustness Grid

For each close-5pp firm-election we compute a cumulative abnormal return over a three-day event window, $CAR_3 = CAR[-1, +1]$, using a Fama–French three-factor adjustment estimated on the $[-130, -11]$ pre-event window. Under the baseline specification used in our pipeline — a uniform mean-difference at $|\text{margin}| \leq 0.05$ with election-clustered standard errors and hedging firms dropped, the uniform convention described in §4 — the pooled 2000–2020 general-election estimate is +17.6 basis points with $p = 0.005$. Under the specification convention Fowler et al. (2020) use to re-examine Akey’s original finding — a local-constant `rdrobust` estimate at a 5-percentage-point bandwidth with firm-clustered standard errors and the same drop-hedgers restriction — the estimate is +29 basis points with $p = 0.14$. Small choices at the specification boundary — kernel weighting (uniform versus triangular), standard-error clustering (election versus firm), and bandwidth-polynomial pairing — move the point estimate in and out of conventional significance without changing its sign.

To place this sensitivity in the broader context of grid-standard robustness, we port the faithful Fowler et al. (2020) Table A8 grid to our panel: four bandwidths $\{1, 2.5, 5, \text{IK}\}$ percentage points crossed with four local-polynomial orders $\{0, 1, 2, 3\}$ on each of two event windows, $CAR_5 = CAR[-1, +5]$ (FGS Panel A) and $CAR_3 = CAR[-1, +1]$ (FGS Panel B). All cells use firm-clustered standard errors and drop double-donor firms, reproducing the FGS Table A8 Stata specification. To calibrate the benchmark, we also port the same grid to Akey’s bundled data (`Akey_RFS_Table_4.dta`, $N = 235$ firm-elections). Table 1 reports summary counts; the full 32-cell output on each sample is in Appendix C.

Two facts organize the grid. First, on our panel 20 of 32 bandwidth–polynomial combinations carry negative point estimates, and one cell, an Akey-style local constant on CAR_3 at a ± 5 pp bandwidth, reaches conventional significance, at +29 basis points with

Table 1: FGS Table A8 robustness grid: our panel vs. Akey benchmark.

Sample	N firm-elec.	Negative / 32	Sig. positive / 32	Sig. negative / 32
Our panel, generals 2000–2020	32,771	20	1	0
Akey (2015) bundled data	235	5	6	0

Note: Cells = 4 bandwidths $\{1, 2.5, 5, \text{IK}\}$ pp $\times$ 4 local-polynomial orders $\{0, 1, 2, 3\} \times 2$ event windows (CAR₅, CAR₃). Firm-clustered standard errors; no fixed effects; double-donor firms dropped. IK bandwidth on Akey’s data uses FGS’s hardcoded values (0.017, 0.0119); on our panel it is computed via CCT-MSERD (`rdrubust`), yielding ≈ 0.006 on the half-margin scale. Significance is at $p < 0.05$. Full 32-cell tables in Appendix C.

$p = 0.034$. One significant cell among 32 tests at the five-percent level is below the 1.6 false positives such a grid produces in expectation under a true null. Nor does the count depend on giving weight to higher-order polynomials: restricting to the specifications closest to modern practice, local constant and local linear at the 2.5pp and 5pp bandwidths, six of those eight cells are positive but only the same single cell is significant. Second, the same grid applied to Akey’s own 235-observation panel yields 5 of 32 negative cells with six significant positives.¹ The instructive comparison is not which sample is more fragile: it is that the benchmark data support the original finding under specification variation while our far larger panel returns an estimate statistically indistinguishable from zero across the standard grid.

We therefore do not claim to replicate Akey’s +3% CAR on donor-firm stock. What our panel contributes on this axis is independent empirical confirmation of Fowler et al. (2020)’s critique of the original finding, on a sample about forty times larger than the thirteen-race special-election panel on which the original magnitude was estimated. Whether close-election resolution moves donor-firm prices at all, under the grid-standard robustness test, is an open empirical question.

¹Our count of negative cells on Akey’s bundled data is similar to but not identical with the 11-of-32 count reported in Fowler et al. (2020); small differences in the specific cell counts across runs are expected given that our sample, event-window coding, and bandwidth-selection pipeline differ from the FGS 2020 setup.

5.2 Subperiod: Pre- and Post-Citizens United

A subperiod split at the January 2010 *Citizens United v. FEC* decision generates a pre-CU null (+5.7 basis points, $p = 0.66$) and a numerically larger post-CU point estimate (+28 basis points, $p < 0.001$ under our baseline election-clustered specification). Given that the pooled estimate fails the FGS grid, we do not pursue this pattern as evidence of an intensifying market reaction. A balance check on pre-period procurement history at the cutoff is clean pre-CU but shows imbalance post-CU, leaving the post-CU subperiod estimate partly composition-contaminated; we return to this in §8. We report the subperiod contrast as a descriptive fact only.

5.3 Long-Horizon Persistence

Setting aside the grid-robustness question, we ask separately whether any initial reaction — however measured — is maintained beyond the event window. If the initial reaction, taken at face value under the baseline specification, correctly forecasts value the connection will subsequently deliver, the price level after the event should not systematically reverse; subsequent abnormal returns should be centered on zero. If the initial reaction encodes an expectation that subsequent information invalidates, the post-event drift may be negative, and the initial magnitude should not be maintained at longer horizons.

We construct a long-horizon abnormal-return panel on the same close-5pp general-election sample. For each of 33,524 firm-elections covering 555 elections in 2000–2020 we compute eight windows of cumulative and buy-and-hold abnormal returns after the event date, plus three pre-event placebo windows, using the same three-factor adjustment on the same $[-130, -11]$ estimation window. The post-event windows begin at trading day +2 to avoid overlap with the three-day event reaction and extend to 252 trading days. Table 2 reports mean differences between winner-supporting and loser-supporting firm-elections, together with cluster-robust p -values and 95% confidence-interval upper

Table 2: Long-horizon CAR persistence, close-5pp generals.

Window	CAR (bps)	SE (bps)	95% CI upper	p -value	Verdict
Event $[-1, +1]$	+17.6	6.3	+29.9	0.005	baseline, sig.
Event $[-2, +2]$	+20.1	7.6	+35.0	0.008	sig.
Event $[-3, +3]$	+7.1	7.9	+22.6	0.37	fades
Drift $[+2, +5]$	-2.6	6.1	+9.3	0.67	≈ 0 , CI rules out +17
Drift $[+2, +21]$	+3.9	13.3	+29.9	0.77	wide CI, no signal
Drift $[+2, +63]$	-53.2	22.0	-10.0	0.016	negative, CI rules out +17
Drift $[+2, +126]$	-40.8	37.9	+33.4	0.28	wide CI, negative point
Drift $[+2, +252]$	-121.0	54.3	-14.6	0.026	negative, CI rules out +17
Placebo $[-63, -11]$	+12.5	17.8	—	0.48	clean
Placebo $[-126, -11]$	+1.1	23.4	—	0.96	clean (BHAR $p = 0.06$)
Placebo $[-252, -11]$	+49.5	36.6	—	0.18	clean

Note: Mean difference in window CAR between winner-supporting and loser-supporting firm-elections at $|\text{margin}| \leq 0.05$, Fama–French three-factor adjusted. Election-clustered standard errors. 95% CI upper bounds are $\hat{\beta} + 1.96 \cdot \text{SE}$. Two counts are involved and are distinct. The long-horizon panel – the set of firm-elections for which the windows are constructed – covers 33,524 close-5pp general-election firm-elections from 555 elections, 2000–2020. The estimation sample – the rows each row of this table is computed on, under the drop-hedgers convention – is 32,082 firm-elections at the event and short-drift windows and declines to 31,378 at the 252-day horizon, since the longest windows require return coverage further past the event. Drift windows begin at trading day +2 to avoid overlap with the event-window reaction. Maintenance of the initial magnitude is statistically ruled out when the 95% CI upper bound falls below +17 basis points. The Event $[-1, +1]$ baseline estimate uses a uniform mean-difference with election-clustered inference under the drop-hedgers convention used throughout (§4); the same sample under the `rdrobust` local-constant convention used by Fowler et al. (2020) (firm-clustered, triangular kernel) yields +29 bps with $p = 0.14$, and the associated bandwidth $\times$ polynomial grid fails conventional robustness (see Table 1).

bounds.

Three features of the table matter. First, the positive event-window point estimate appears at the two narrowest nested windows: $\text{CAR}[-1, +1]$ and $\text{CAR}[-2, +2]$ are similar and significant under election-clustered inference, while $\text{CAR}[-3, +3]$ has already faded to an insignificant +7.1 basis points. The reaction is confined to a narrow window around the election and begins to decay within the same week. Second, across the five post-event drift windows not one point estimate is significantly positive. Three of the five – drift $[+2, +5]$, drift $[+2, +63]$, and drift $[+2, +252]$ – have 95% confidence-interval upper bounds strictly below +17 basis points, statistically ruling out that the initial event-window magnitude is maintained at those horizons. Third, the pre-event placebo windows are mostly clean under CAR; the one exception is the pre-126-day buy-and-hold

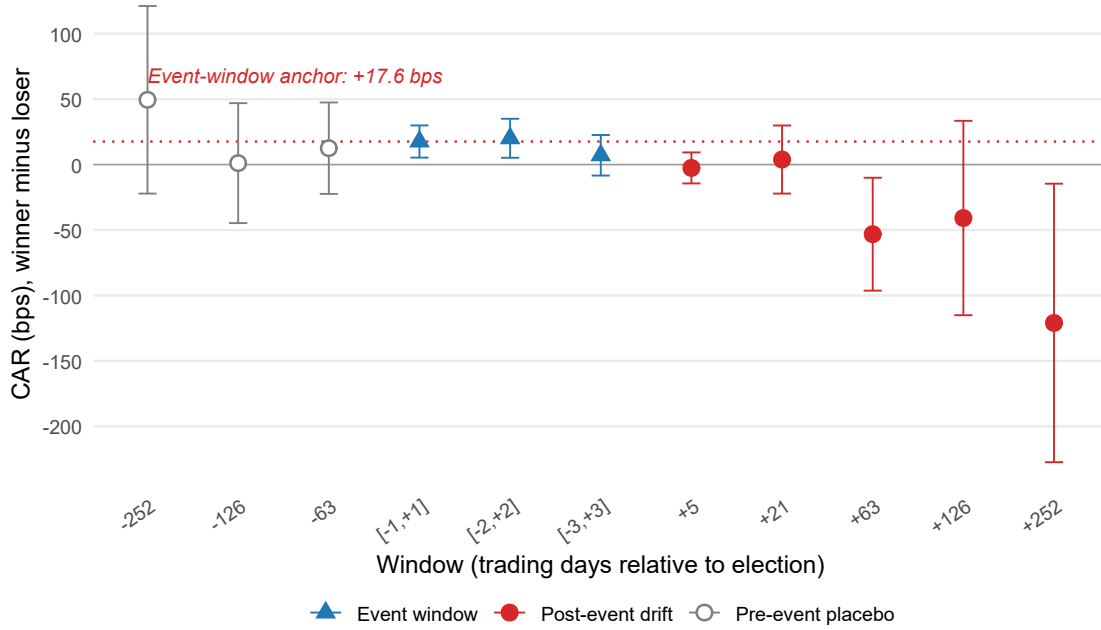


Figure 2: CAR persistence on winner-supporting firms, close-5pp generals 2000–2020. Points are mean-difference estimates (winner minus loser) of cumulative abnormal return over the indicated window, election-clustered 95% CIs shown. The event-window reference line (+17.6 bps, dotted) marks the baseline-specification point estimate whose cross-specification robustness is examined in Table 1; post-event drift windows are centered on or below zero, with the $[-2, +63]$ and $[-2, +252]$ CIs strictly below that reference line. Full horizon table above.

placebo, which comes in at -22.5 basis points with $p = 0.057$ – a mild warning about the long-horizon BHAR noise that Barber and Lyon (1997) and Kothari and Warner (2007) have flagged in related settings.

What these numbers defensibly support is a claim of no positive persistence. At no post-event horizon does the evidence favor maintenance of the initial reaction, and at three of five horizons the evidence actively rules it out at the 95% level. What these numbers do not defensibly support is a claim of reversal. The pattern is not monotonically decreasing, the long-horizon buy-and-hold placebo is marginally anomalous, and multi-window testing across sixteen post-event estimates with a Bonferroni-corrected threshold of $p < 0.003$ would reject only the initial event window. We therefore adopt the more conservative no-positive-persistence reading: across all post-event horizons the evidence is consistent with the initial reaction being at least unwound, but we do not claim a full

reversal. Even taken at face value under the baseline specification, the initial reaction is not maintained as a positive return over the subsequent weeks and months.

5.4 Credit-Market Reaction: Bond Spreads

As a credit-market analog to the equity event-window, we apply the uniform Δ -RDD to the issuance-amount-weighted yield-to-maturity spread relative to comparable-duration Treasuries (described in §3). The pooled Δ -RDD returns -0.00024 with standard error 0.00021 and $p = 0.25$, estimated on 24,586 close-5pp firm-elections from 541 close elections; the underlying spread panel covers 12,945 firm-years on 969 public-debt-issuing firms, 61% of the main sample. The sign is consistent with winner-firm spread compression, but the estimate is indistinguishable from zero. At approximately 2.4 annualized basis points the magnitude is small in absolute terms, and the 95% confidence interval rules out compression larger than approximately 6.5 basis points – below the credit-market magnitudes the political-connections literature has documented on cost-of-capital and financing-choice outcomes. The pre-trend placebo is clean at $p = 0.68$. Subperiod estimates are separately null. As a check on whether credit-market effects emerge only in crisis states – where government support might selectively compress connected-firm spreads – we add a $\text{win} \times \text{crisis_post}$ interaction under the same two crisis definitions used in Appendix E; both interaction estimates are null. The credit market, like the equity market beyond the event window, does not price a sustained connection premium. Whether any component of the brief equity reaction corresponds to realized allocation decisions on the firm’s balance sheet or in the government’s own records is the question we take up in §6.

Table 3: Government-benefit outcomes, Δ -RDD pooled 2000–2022.

Outcome	$\hat{\beta}$	SE	p -value	Placebo p	Verdict
Procurement ($\Delta \log(1 + \cdot)$)	+0.006	0.045	0.88	0.26	null
ETR ($\Delta \text{etr_gaap}$)	+0.0012	0.0024	0.63	0.61	null
AAER violation ($\Delta I(\text{any})$)	−0.0010	0.0021	0.62	0.95	null, sign as hypothesized
Grants, total assistance	−0.198	0.073	0.007	0.0005	placebo fails (composition)
Grants, any assistance	−0.022	0.007	0.001	0.0003	placebo fails (composition)
Grants only (excl. loans)	+0.053	0.026	0.041	0.13	borderline

Note: Δ -RDD estimates on close-5pp firm-elections, cluster-robust (election) standard errors. Pre window = $\{t - 1, t\}$, post window = $\{t + 1, t + 2\}$, placebo earlier-pre window = $\{t - 3, t - 2\}$. 2024 cycle excluded due to truncated post-window. Hedging firms excluded. Each row is the coefficient on the win indicator in equation (2). Placebo column reports p -value from the analogous regression on $\Delta y^{\text{placebo}}$ (equation (3)); placebo $p < 0.05$ signals pre-trend contamination.

6 Government Benefits to Donor Firms

The connections-matter literature’s implicit claim is that a tie to a winning legislator produces benefits government can actually deliver: contracts, tax treatment, enforcement posture, grant dollars. This section tests that delivery claim directly, on the government’s own administrative records and the firm’s own filings, without conditioning on whether the event-window reaction examined in §5 is accepted. The four outcomes we test – procurement allocations, effective tax rates, SEC enforcement exposure, and federal grants – cover three substantive channels through which government can transfer resources to a firm: direct contracting, tax treatment, and regulatory posture. For each outcome we apply the uniform Δ -RDD specification developed in §4: differencing within firm-election, requiring a clean pre-trend placebo, and triangulating across clustered OLS, `rdrobust`, and `rdrandinf`. Table 3 summarizes the pooled results.

6.1 Procurement

Procurement is the outcome with the strongest prior expectation in the literature (Agca and Igan, 2023; Baltrunaite, 2020; Goldman et al., 2013) and the one on which the level-RDD composition bias is most acute. The pooled effect on log-procurement is a precise

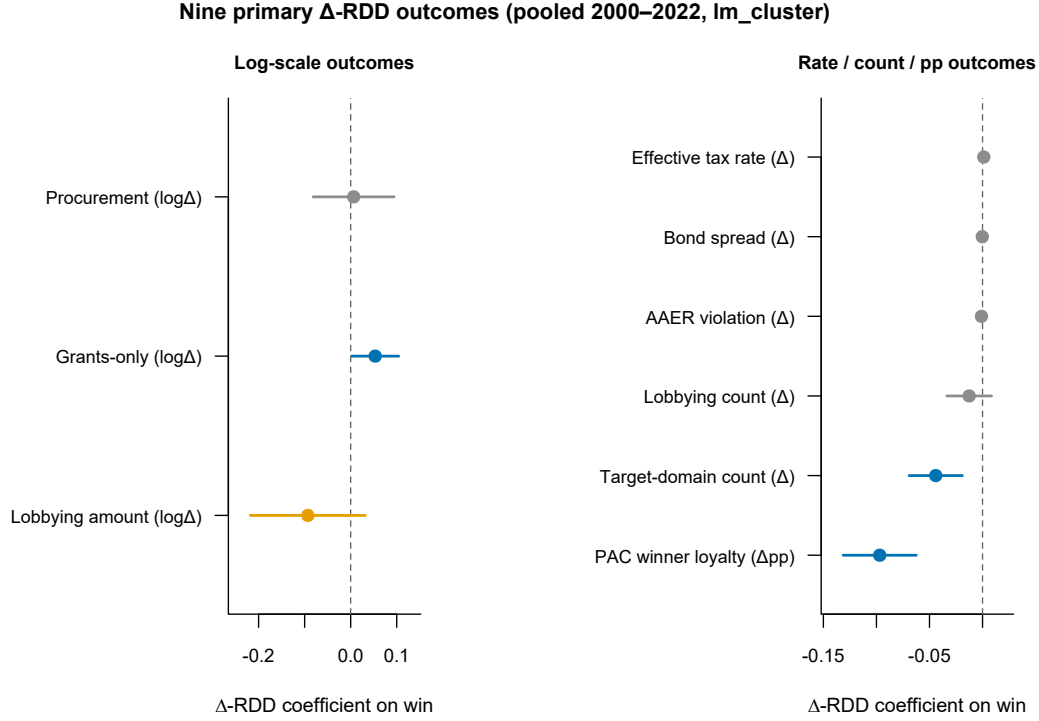


Figure 3: Nine primary coefficients on the win indicator, pooled 2000–2022, with election-clustered 95% confidence intervals. Left panel: three log-scale outcomes. Right panel: six rate-, count-, and percentage-point-scale outcomes. Grey markers are null; amber marks the borderline lobbying-expenditure coefficient ($p = 0.14$); blue marks the three statistically significant coefficients (grants-only, lobbying target-count, PAC winner-party share). Each significant coefficient is small and qualified in the main text: in particular, the PAC winner-party share has a failed placebo that identifies the pattern as pre-period composition rather than causal response (§7.3).

zero – a point estimate of $+0.006$ log-units ($SE = 0.045$; $p = 0.88$), on an estimation sample of 35,033 firm-elections from 588 close elections. The pre-trend placebo is null, at $+0.073$ log-units ($p = 0.26$). A 95% confidence interval of approximately $[-0.08, +0.09]$ on the differenced log-procurement outcome rules out effects larger than about ten percent on the log scale – tighter than the level-RDD coefficient we ourselves recover on the same panel before differencing ($+0.308$, §4), and below the defense-contract allocation advantage documented in Agca and Igan (2023) and the ban-induced contraction in donors’ winning probability documented by Baltrunaite (2020). The channel that carries the strongest prior for delivery returns a tightly estimated zero on candidate-level electoral variation.

Splitting by Citizens United yields no winner premium in either subperiod, though the two point estimates fall on opposite sides of zero. In the pre-CU 2000–2008 subperiod, the OLS estimate is +0.139 log-units ($p = 0.030$) with a clean placebo (+0.025, $p = 0.71$), which in isolation would support a small positive winner premium; but estimator triangulation fails, with `rdrobust` at +0.10 ($p = 0.61$) and `rdrandinf` at +0.30 ($p = 0.16$), both null, and under the rule we adopt in §4 a single OLS-significant estimate unechoed by local-polynomial or Fisherian inference does not support a causal claim. In the post-CU 2010–2022 subperiod the main Δ -RDD estimate is -0.089 log-units ($p = 0.10$) on 20,997 firm-elections, with a clean pre-trend placebo (+0.107, $p = 0.15$). The pooled null therefore does not depend on either subperiod alone: each subperiod is separately null with a clean placebo, the within-firm fixed-effects specification is null in every subperiod with a pooled 95% confidence interval of $[-0.05, +0.05]$ (§8), and the one-close-race slice of §4.4 returns the same verdict on the firm-cycles where cross-race aggregation cannot operate.

Our finding joins the growing null-findings literature on corporate campaign-contribution returns: Fowler et al. (2020)’s firm-level stock-return null and Fournaies and Fowler (2021)’s industry-level policy null, together with our candidate-level close-election RDD, document nulls through three different designs on three different realized-outcome margins.

6.2 Other Government-Benefit Channels

The second and third realized-benefit channels – effective tax rates and SEC enforcement exposure – show the same pattern: null point estimates with tight confidence intervals that rule out the canonical literature effect sizes. These two estimates carry a stronger evidential status than the procurement null. For procurement the differencing had bias to remove: the level design was composition-contaminated, and the Δ -RDD repaired it. For effective tax rates and enforcement there is no composition to remove: the level regres-

sions themselves are null on both the pre and the post window (Appendix B), so the cutoff assigns treatment against balanced baselines and differencing serves only to tighten precision. Together with the bond-spread estimate of §5.4, whose levels are likewise balanced, these are the cleanest causal estimates in the paper, and they are nulls. The pooled effect on the GAAP effective tax rate is +0.12 percentage points ($SE = 0.24$; $p = 0.63$), estimated on 31,703 firm-elections from 586 close elections; the 95% confidence interval of approximately $[-0.004, +0.006]$ rules out responses larger than about half a percentage point, roughly a quarter of the smallest effects the corporate-tax-connections literature documents. The pooled effect on the any-AAER-violation indicator is -0.10 percentage points ($SE = 0.21$; $p = 0.62$) on 35,611 firm-elections from 588 close elections: sign-consistent with a winning-reduces-enforcement-exposure hypothesis, but indistinguishable from zero on a 2–3% base-rate outcome, with a 95% confidence interval that rules out relative effects on violation probability of roughly 20% or more. For both outcomes the pre-trend placebo is clean and the estimator triangulation agrees on null. Subperiod estimates are separately null. A release-sensitivity specification on AAER is null under OLS in every subperiod, and its post-CU cell is a three-frame disagreement: Fisherian `rdrandinf` returns a significant estimate of the opposite sign to OLS, and `rdrobust` on this rare binary is degenerate. Under our triangulation rule that is a cell where the design does not identify, so we document it as a footnote-level caveat and do not read it as evidence of enforcement leniency. Detailed subsection-level treatment of each outcome, including robustness across subperiods and placebo tables, appears in Appendix D. These nulls, together with the procurement null, join Fowler et al. (2020) and Fourirnaies and Fowler (2021) as documentation that corporate campaign activity does not produce measurable realized returns on the outcome margins adjacent literature identifies as delivery-capable.

6.3 Federal Grants and Assistance

Federal non-procurement outlays from USASpending FABS comprise grants, loans, loan guarantees, direct payments, and insurance programs. We report three variants because the aggregate is heterogeneous in its political economy. The total-assistance estimate is -0.198 log-units ($p = 0.007$): a strongly significant negative coefficient at face value. The placebo, however, fails loudly at $+0.179$ with $p = 0.0005$, identifying the pattern as a composition trajectory rather than a causal effect. The any-assistance indicator shows the same structure. The grants-only variant – restricted to the FABS assistance types unambiguously classified as grants, excluding loans, loan guarantees, direct payments, and insurance – behaves differently. The pooled estimate is $+0.053$ log-units ($SE = 0.026$; $p = 0.041$), about five percent; the placebo is borderline, at $+0.047$ log-units ($p = 0.13$). Two considerations keep us from reading this cell as delivery. The placebo point estimate is nearly as large as the treatment estimate, so much of the same-signed movement predates the election; and the cell is one of nine primary outcomes, a family in which nine tests at the five-percent level produce 0.45 false positives in expectation, so a single $p = 0.041$ coefficient carries little evidential weight on its own. Grants-only is nonetheless the substantively cleanest slice of federal assistance – direct budget allocation where appropriations-committee discretion plausibly operates – so if delivery happens anywhere in our benefit family it should happen here, and what appears here is a small estimate that does not separate from its own pre-trend. The paper’s thesis does not depend on reading it either way.

6.4 Statistical Power and Minimum Detectable Effects

Before interpreting the null pattern as informative, we verify that the design is adequately powered to reject the effect magnitudes the adjacent literature has documented. A null is an argument only when the design could have detected the effects the literature

Table 4: Minimum detectable effects, Δ -RDD pooled 2000–2022.

Outcome	Scale	SE	MDE	MDE in relative terms
Procurement	$\Delta \log(1 + \$)$	0.0448	0.125	13% of obligated dollars
Grants only (excl. loans)	$\Delta \log(1 + \$)$	0.0261	0.073	7.6% of assistance dollars
Lobbying expenditure	$\Delta \log(1 + \$)$	0.0637	0.178	20% of lobbying dollars
Lobbying reports	$\Delta \log(1 + n)$	0.0108	0.030	3.1% of report count
Target-domain reports	$\Delta \log(1 + n)$	0.0128	0.036	3.7% of target-code report count
Effective tax rate	Δ rate	0.0024	0.0068	0.68 pp of pre-tax income
Bond spread	Δ rate	0.00021	0.00058	5.8 basis points
AAER violation	$\Delta I(\text{any})$	0.0021	0.0059	20–30% of a 2–3% base rate
PAC winner-party share	Δ share	0.0176	0.049	4.9 pp of winner-party dollar share

Note: $MDE = 2.8 \times \widehat{SE}$ is the two-sided minimum detectable effect at the five-percent level with 80% power. Standard errors are the pooled Δ -RDD election-clustered standard errors reported in Table 3 and Table 5. For log-scale outcomes the relative term is $\exp(MDE) - 1$; for rate outcomes it is the bound expressed in the units the adjacent literature uses, and for the AAER indicator it is the bound divided by the 2–3% within-window base rate.

claims (Rainey, 2014), and equivalence-style reporting for regression-discontinuity settings formalizes the same demand (Hartman, 2021); the minimum-detectable-effect and confidence-interval bounds we report throughout implement it. For a two-sided test at the five-percent level with 80% power, the minimum detectable effect on the Δ -RDD specification is $MDE = 2.8 \times \widehat{SE}$. Table 4 applies that formula to all nine primary outcomes and converts each bound into the units in which the relevant literature reports its estimates.

For procurement, an MDE of 0.125 log-units – about a thirteen percent change in obligated dollars – lies below the +0.28-order (log-unit) allocation advantage in Agca and Igan (2023)’s defense-contract setting and below our own panel’s level-RDD coefficient of +0.308 (\$4); Baltrunaite (2020)’s Lithuanian ban documents a 5-percentage-point reduction in donors’ winning probability, a scale our design is also powered to rule out on the implied log-dollar margin. For effective tax rates, the MDE of 0.68 percentage points is between one-third and one-seventh of the 2–5 percentage point effects documented in the corporate-tax-connections literature. The bond-spread MDE is 5.8 basis points, below the 20–50 basis-point compression the credit-market connections literature reports.

Enforcement is the one channel whose bound does not clear the literature’s range with room to spare, and the comparison is interpretable only in relative units because Correia

(2014) reports relative rather than absolute reductions. Dividing the AAER violation MDE of 0.59 percentage points by the within-window violation rate of 2–3 percent gives a detectable relative reduction of 20 percent at a 3 percent base rate and 30 percent at a 2 percent base rate. Correia (2014)’s SEC-enforcement specifications report relative reductions of 33–60 percent; at a 2 percent base rate the bottom of that range is 0.66 percentage points, only marginally above our MDE. Our design therefore has 80% power against the upper half of the enforcement-leniency range and almost no margin against its lower end, and we do not claim to rule out the smallest reductions Correia (2014) reports.

The single non-null in the benefit family – grants-only at $+0.053$ log-units, $p = 0.041$, borderline placebo at $p = 0.13$ – is itself within the detection range and is quantitatively small. Subject to the enforcement qualification above, the benefit-family nulls are informative, not underpowered.

6.5 Consolidating the Four Outcomes

Across procurement allocations, effective tax rates, SEC enforcement actions, and federal grants, the realized-benefit channel shows a coherent pattern. Three of four outcomes return clean nulls whose 95% confidence intervals rule out the magnitudes the adjacent literature has documented. The fourth, grants-only, is a small positive that does not separate from its own placebo and is within the family’s false-positive expectation. The magnitude of realized government-sourced benefit is small where it exists at all and invisible where it does not. A referee concern on the null pattern is that delivery may operate only off the equilibrium path – visible in crisis states rather than in average times. Our design captures this possibility via crisis-state interaction specifications on five outcomes under two crisis definitions – a tight definition covering the financial-crisis and COVID cycles and a broad definition that adds the dot-com and 9/11 cycles; across the resulting ten interaction tests, no coefficient reaches conventional significance ($p \in [0.21, 0.98]$; see Appendix E). The off-equilibrium reading does not rescue realized delivery on our design.

Table 5: Political-behavior outcomes, Δ -RDD pooled 2000–2022.

Outcome	$\hat{\beta}$	SE	p -value	Placebo p	Verdict
Lobbying $\Delta \log(1 + \$)$	−0.093	0.063	0.14	clean	null, slight negative
Lobbying $\Delta I(\text{any})$	−0.006	0.004	0.17	clean	null
Lobbying $\Delta \log(1 + n_{\text{reports}})$	−0.013	0.011	0.24	clean	null
Target count $\Delta \log(1 + \cdot)$	−0.044	0.013	0.0006	0.09	small negative, plac. marginal
Target share Δ	−0.001	0.002	0.60	0.46	null
PAC $\Delta \log(1 + \$_{\text{cycle}})$	+0.43	0.27	0.12	0.31	null, direction +
PAC $\Delta \log(1 + n_{\text{cands}})$	+0.16	0.087	0.073	0.22	marginal +
PAC $\Delta \log(1 + \$_{\text{opp. party}})$	+0.44	0.29	0.13	0.99	null, direction +
PAC Δ winner-party share	−0.097	0.027	< 0.001	0.0003	placebo fails
PAC follow-up loyalty (post)	+0.228	0.021	< 0.001	(+0.292 pre)	taper, not amplification

Note: Δ -RDD estimates, close-5pp firm-elections, cluster-robust (election) standard errors. Lobbying and target-alignment outcomes use fiscal-year pre = $\{t - 1, t\}$, post = $\{t + 1, t + 2\}$. PAC outcomes use cycle-based windows (pre = $t - 2$, post = $t + 2$, placebo = $t - 4$). Follow-up loyalty row reports the post-cycle mean difference in the probability of donating to the close-race winner, between firms that had donated to him before the election and firms that had donated to his opponent; the pre-cycle value in parentheses is the same difference measured two cycles before the election.

Whether firms themselves behave as if the connection delivers, as a revealed-preference check on our measurement, is taken up in §7.

7 What Donor Firms Do Next: A Behavioral Consistency Check

The null results in §6 describe what donor firms receive from government after a close-election win. This section asks what donor firms *do* after the same event. The revealed-preference logic developed in §2 is that firms that have just acquired a politically valuable connection should subsequently invest more in exploiting it: more lobbying expenditure, more lobbying targeted at the connected legislator’s committee jurisdictions, and PAC-contribution portfolios rebalanced toward the newly connected insider. We implement the uniform Δ -RDD of §4 on three outcome families – lobbying expenditure, lobbying target alignment, and PAC contribution strategy. Table 5 summarizes the pooled estimates.

7.1 Lobbying Expenditure

We observe federal lobbying through the LobbyView parse of LDA quarterly filings, aggregated to the `GVKEY`-fiscal-year level, with zero-filling of non-appearing firm-years. Three outcome variants cover different margins of adjustment: log-dollar expenditure, the any-lobbying indicator, and the log number of reports. All three pooled estimates are null under clustered OLS: -0.093 ($p = 0.14$) for log-dollar, -0.006 ($p = 0.17$) for the indicator, -0.013 ($p = 0.24$) for log-count. All three placebos are clean. The 95% confidence interval on the log-dollar variant, $[-0.22, +0.03]$, rules out post-election lobbying increases larger than about three percent on the log scale – a bound tight enough to rule out the effect sizes that a patronage reading of the Hall and Wayman (1990) committee-access theory or of Bertrand et al. (2014)’s lobbyist-connections framework would predict for a newly formed committee connection. The `rdrobust` estimators return positive coefficients of $+0.20$ to $+0.40$, but their placebos are also positive and significant (pooled placebo $p = 0.04$), indicating that the local-polynomial frame is picking up composition in the ever-lobbying indicator that the differencing absorbs under OLS. We rely on the frame whose placebo is clean; OLS is the trustworthy frame here, and OLS is uniformly null.

All three OLS point estimates are negative, though none is individually significant, and the negative direction concentrates in the post-CU subperiod, where the log-dollar estimate is -0.178 ($p = 0.080$) and the Fisherian frame returns -0.216 ($p = 0.086$) (Table 11). The cell reaches conventional significance under neither frame, and we read it as we read the pooled estimate: a null with a negative sign, not a finding. A negative direction would be consistent with a direct-access substitution story in which a firm that has acquired a friendly insider reduces intermediate lobbyist expenditure. We are underpowered to distinguish substitution from a pure null and do not endorse the substitution reading. We do note that the direction is inconsistent with a patronage reading on which a new connection has made marginal lobbying dollars more productive and hence more

worth spending.

7.2 Lobbying Target Alignment

Lobbying expenditure is untargeted. A sharper test is whether firms shift their lobbying composition toward the issue codes associated with the newly connected legislator’s committee jurisdictions. We construct this test by recovering the winner’s `bioguide_id` and `icpsr` via a join to Voteview HSall, attaching the winner’s committee assignments from Charles Stewart’s House Committee Assignments file, and applying the eighteen-committee-to-LDA-issue-code crosswalk documented in Appendix A. Two outcomes follow: a count outcome, $\Delta \log(1 + \text{reports on target codes})$, and a share outcome, the change in the fraction of the firm’s lobbying reports on target codes.

The count Δ -RDD returns -0.044 with $p = 0.0006$: a small but statistically significant negative coefficient; the placebo is marginal at $p = 0.09$. Winner-supporting firms lobby slightly less on the newly connected legislator’s committee-domain issue codes in the post-election window than loser-supporting firms do; the magnitude is small (roughly four percent on the log scale) and the placebo marginality tempers the causal claim, but the direction is opposite to what a patronage reading would predict. The share Δ -RDD returns -0.001 with $p = 0.60$: null. What the level regressions reveal is substantively interesting. The pre-election share and the post-election share each regress onto the win indicator with a coefficient of $+0.10$ and $p < 10^{-14}$: firms donate to legislators whose committee jurisdictions already match the issue codes they are already lobbying on, and this near-mechanical selection is visible before the election happens. After differencing, the pre-existing composition cancels, and the Δ -RDD coefficient on share is zero. The election outcome itself does not rebalance the composition of lobbying. Taken together, the target-count and target-share results are inconsistent with a patronage reading of the connection: if the newly connected legislator’s committee jurisdictions were a more valuable lobbying target after the connection formed, firms should be lobbying those jurisdictions

more in absolute terms or in relative terms. Neither moves up; count moves slightly down.

7.3 PAC Strategy: Strategic Giving, Unrealized Returns

Corporate campaign contributors behave strategically. They give more to the legislators sitting on committees that regulate them (Fouirnaies and Hall, 2018), and they pay a financial incumbency premium on the seats their industry depends on (Fouirnaies and Hall, 2014). PAC giving is shaped more by access considerations than by ideology (Barber, 2016). All of this behavior is consistent, as Fouirnaies and Fowler (2021) put it for insurance firms, with the view that firms act as if they are getting something in return. Yet in the close-election slice where a strategic investment most plausibly pays off – the donor-supported candidate narrowly wins – firms neither concentrate subsequent giving on the winner nor lobby more on his committee’s jurisdiction. Our PAC-strategy evidence documents this strategic-versus-realized paradox on the firm’s own revealed-preference margin.

A firm’s PAC portfolio adjusts at the cycle level, so we use cycle-based windows: pre is cycle $t-2$, post is cycle $t+2$, placebo is cycle $t-4$. The pooled estimate on total PAC dollars is $+0.43$ log-units ($p = 0.12$) – directionally positive but not significant, placebo clean. Distinct recipient candidates return $+0.16$ log-units with $p = 0.073$: marginally positive with a clean placebo. Winner-side firms broaden their donor network slightly after the election, giving to a wider set of candidates than loser-side firms do. Winner-party share shows -0.097 with $p < 0.001$ under OLS but a strongly positive placebo ($+0.052$, $p = 0.0003$): trajectory-reversal composition that we do not interpret substantively. The sharp loyalty test asks: how much more likely is a firm that had donated to the close-race winner to donate to him again in the next cycle, compared to a firm that had donated to his opponent? In the pre-cycle this difference is 29.2 percentage points – near-mechanical, since pre-cycle donation to that candidate is what defined a firm as winner-side in the

first place. In the post-cycle the same difference falls to 22.8 percentage points. Winner-side firms concentrate less on the newly connected legislator, not more. The loyalty gap tapers by 6.4 percentage points rather than amplifying.

Where the additional dollars go is itself informative. Decomposing the PAC budget by party, winner-side firms' giving to the party that just lost the seat rises by +0.44 log-units ($p = 0.13$) against a placebo of +0.002 ($p = 0.99$); descriptively, winner-side firms add about \$26,000 per cycle in opposition-party giving, against about \$8,000 for loser-side firms. The estimate is not individually significant, but its placebo is clean and its direction is the one the portfolio account of §2 predicts: after a win, firms extend giving toward the other party rather than concentrating on the winning side. The complementary decomposition, giving to the winning party's candidates, returns a near-zero estimate whose placebo fails (+0.03, $p = 0.90$; placebo $p = 0.009$), so we do not interpret that cell.

Taken together, the PAC results match the portfolio account of §2 on every measured margin: winner-side firms expand their candidate network (number of distinct recipients marginally up), disperse their contributions relative to the specific winner (follow-up loyalty down), extend giving toward the opposition party, and keep total expenditure approximately flat. The patronage account predicts the opposite on each margin: total expenditure up, candidate count flat or down, giving concentrated on the winning side, and specific-winner loyalty sharply up. Three interpretations of the strategic-versus-realized paradox remain available. First, consistent with Fourniaies and Fowler (2021), firms may find it as difficult to estimate the realized returns to their campaign contributions as researchers do, and their strategic-giving behavior may operate under genuine inferential uncertainty about what the contributions actually buy. Second, a shareholder-versus-government-relations agency problem may permit persistent strategic giving even without verifiable realized returns. Third, our preferred complement to these two readings, the portfolio account itself: winner-side firms rebalance broadly across an expanded candidate set rather than doubling down on the specific winner, the behavior of a portfolio

manager rather than of a client maintaining a binding relationship with an insider.

7.4 Consolidating the Three Behaviors

Across lobbying expenditure, lobbying target alignment, and PAC portfolio adjustment, firms do not act as if a close-election win has made their connection more valuable. Firms, who have more information about the productivity of their own political connections than outside investors have at the moment of an election result, do not respond to a close-election win by increasing their investment in the connection. The behavioral response is the revealed-preference signal that, whatever the stock market prices at the moment the race is called, firms themselves do not treat a marginal win as a substantive upgrade in leverage worth doubling down on. The measurement-side null on realized delivery and the behavior-side null on investment are mutually reinforcing, and together they make it difficult to sustain a reading on which the connection genuinely delivers benefit that is merely below our outcome-data resolution.

8 Robustness

We report four robustness checks that together close the main causal-identification loop across our Δ -RDD outcomes; detail is in Appendix F. First, within-firm fixed effects on the 1,047 firms that appear on both winner and loser sides of different close races return a pooled procurement coefficient of -0.002 ($p = 0.92$), a 95% confidence interval of $[-0.052, +0.047]$, and subperiod coefficients that are individually null. Using only the dual-appearance subpopulation, and absorbing firm-specific level via fixed effects rather than differencing, the within-firm FE specification converges to the same tightly estimated zero on pooled data. Second, augmenting the general-election panel with 67 close House special elections (2002–2024) attenuates rather than amplifies the naive level-RDD coefficient: the general-only estimate of $+0.153$ falls to $+0.112$ under the augmented panel, and

the $\text{win} \times \text{is_special}$ interaction is directionally negative with $p = 0.13$. A pipeline-error or missing-identifying-variation reading would predict amplification. The attenuation pattern is consistent with a general across-subtype null. Third, the Citizens United split partitions the pooled sample into 2000–2008 and 2010–2022, and the pooled null is not an artifact of pooling across that regime change: each subperiod’s procurement Δ -RDD is individually null with a clean pre-trend placebo (§6.1). The split matters because the donor pool changes character across it: the close-race hedging rate fell from 4.3% in 2010 to 0.6% in 2022, so the later cycles compare increasingly committed single-side donors. That evolution appears as the level imbalance the two subperiods share, not as differential trend imbalance that would defeat the differenced design in one regime. We treat the hedging collapse as a descriptive fact about the donor pool rather than as a causal claim about the decision itself; Citizens United exposure is universal across firms and election cycles, so a clean cross-sectional counterfactual does not exist within our design. Fourth, stratifying the Δ -RDD by firm portfolio breadth rejects cross-race dilution as a source of the null. A reader concerned that firm-level outcomes aggregate across a diversified portfolio of legislator connections – so any single close-race win/loss is mechanically attenuated at the firm level regardless of whether connections deliver per race – would predict that restricting to single-candidate firm-cycles recovers a positive delivery effect and that the $\text{win} \times \log(n_{\text{cands}})$ interaction is negative and significant. Neither prediction survives the data. Across all nine primary outcomes the interaction coefficient is either null or signed opposite to the dilution hypothesis, and single-candidate subsample estimates are null or negative for every outcome; detail in Appendix F.5, with the complementary one-close-race-per-firm-cycle slice, which removes cross-race mixing by construction on a nine-times-larger subsample, in §4.4 and Appendix F.6. Alternative sample restrictions – retaining hedging firms and assigning them to the winner or loser side by the sign of their net donation, tighter bandwidth at $|\text{margin}| \leq 0.03$, inclusion of the 2024 cycle – do not change the sign or significance of any of our primary outcomes.

9 Discussion

9.1 Prices Without Rents

A natural reading of the close-election literature’s cumulative-abnormal-return evidence is that positive event-window reactions, where they are found, reflect information updating about legislator identity and future committee composition rather than anticipation of realized rents that subsequently flow to donor firms. The narrative tradition running from Goldman et al. (2009) through Cohen et al. (2013) and Akey (2015) rests on an inferential step that every link in the chain defends except the last: efficient markets do price political events; they price them into stock prices; the reactions are reported as positive and in the expected direction. The robustness of those reactions to standard specification grids has itself been questioned (Fowler et al., 2020) and, on our sample, does not survive the same grid (§5). What the chain does not support is the identification of the *expected* benefit the market is pricing with the *realized* benefit that subsequently materializes. A positive cumulative abnormal return is a valid estimator of market expectation. It is an invalid estimator of realized delivery unless one is prepared to assume ex ante that expectations are unbiased on the specific unobserved channel that delivers benefit – an assumption neither logically required nor, in the data we have assembled, empirically supported. The sharpened reading we propose is that event-window reactions in close-election CAR designs, where and when they arise, are best interpreted as information pricing: the market updates on the identity of the winning legislator and on whatever expectations flow from that identity, but the update does not correspond to a realized benefit that subsequently flows from the government to the firm. Connections, when they move prices at all, do so only momentarily; they do not move rents.

9.2 Why No Realized Delivery?

The productive question our evidence raises is why, in the setting the close-election literature has identified as the most likely case for connection-value effects, realized fundamentals do not move and firms themselves do not behave as if they have acquired valuable leverage. Four alternative explanations organize the debate, and our evidence bears differently on each. First, and our primary reading: to the extent the event-window reaction is accepted at face value under its baseline specification, the cumulative abnormal return is an information-pricing event. Markets aggregate private information about legislator identity and future committee composition; no realized-rent forecast is priced, and the absence of subsequent realized delivery on fundamentals or in firm behavior is the expected implication. That any positive initial reaction does not persist beyond the event window (§5) is internally consistent with this reading: to the extent the market briefly updates, it revises as subsequent information arrives, and the initial magnitude is not maintained. Second, consistent with Fourniaies and Fowler (2021), firms themselves may find it as difficult to estimate returns to political contributions as researchers do, and strategic-giving behavior may operate under genuine inferential uncertainty about what the contributions actually buy. Under this reading the paradox between documented strategic giving and null realized returns is a paradox of limited inference, not of misaligned incentives. Third, delivery may operate through channels below our detection floor: long-horizon goodwill accumulated beyond our two-year window, narrow programs (TARP Capital Purchase Program, COVID sub-program emergency procurement, earmarks) captured only below FPDS and FABS aggregation, or sub-contract discretion not indexed by our data. Our crisis-interaction tests (Appendix E) rule out broad-stroke crisis-state delivery on our four measured outcomes, but narrow-channel delivery remains empirically open. Fourth, firm-specific strong-tie returns – private concessions to individual firms that our design averages over – could in principle operate. Our top-decile-connection-strength appendix analysis returns nulls of comparable magnitude (Appendix F), which constrains

but does not rule out this reading. Our paper does not resolve among the four; it establishes the puzzle’s empirical grounding and identifies information pricing as the simplest reading compatible with the full pattern.

A recent wave of positive findings on connection returns does not contradict the null, because none of it estimates our contrast. Denes and Scanlon (2025) find that firms channeling dark money receive more procurement in observational comparisons; Fotak et al. (2025) document that politically connected firms obtained more exemptions from the 2018 tariffs, and Houston et al. (2026) that locally connected firms win more federal contracts; Wu and Cao (2026) show that product-market peers respond to a donor firm’s special-election win by appointing politicians to their own boards. The observational designs compare connected with unconnected firms, a margin on which our own level regressions also return large positive coefficients before the pre-period placebo attributes them to selection; the peer-response result measures what firms believe and how they govern, not what government delivers. The one directly comparable design, Bisetti et al. (2026)’s close-election RDD showing that a Democratic win reduces pollution at local plants, differs on exactly the margin that breaks ours: their treatment is the elected legislator’s party in the district where a plant sits, an assignment the firm does not choose, whereas our treatment status derives from the firm’s own decisions about which races to fund. Donor-side selection is the mechanism our composition diagnostics identify, and it is absent from their assignment by construction, which is why a firm-level close-election RDD can be clean there and compromised here. Read together, the recent results are consistent with our conclusion while narrowing its scope: correlations between connections and benefits are abundant, and what fails to appear is delivery identified off the donation-defined close-election contrast.

9.3 Portfolio Versus Patronage

A second substantive lens the evidence invites concerns the shape of corporate political strategy. Section 2 set out the two candidate accounts and their competing predictions: the patronage account, rooted in the political-science tradition on clientelism, the finance tradition on revolving doors, and the public-choice tradition on rent-seeking, treats political ties as relationship-specific assets that reward post-win concentration; the portfolio account treats ties as partially substitutable access instruments that reward post-win rebalancing. The political-behavior results in §7 separate the accounts in the portfolio direction on every measured margin: winner-side firms slightly broaden their candidate roster; their donation probability to the legislator they had backed exceeds the loser-side baseline by 23 percentage points in the post-cycle, down from 29 points before the election, tapering where patronage predicts amplification; lobbying expenditure is flat or slightly down; lobbying composition does not shift toward the newly connected legislator's committee jurisdictions; and the one cleanly identified movement in PAC dollars is an expansion of opposition-party giving. We offer the portfolio account as a complementary third reading alongside the inferential-uncertainty and shareholder-agency-problem interpretations that Fourinaies and Fowler (2021) flag. Corporate political activity in U.S. federal elections looks more like exposure management across substitutable relationships than like personalistic clientage.

9.4 Policy Implications

Whatever positive market reaction exists in this setting, our evidence indicates that it does not correspond to realized rents flowing from the government to the donor firm. That is a different normative question from the quid-pro-quo concern motivating most campaign-finance regulation. Meaningful efforts to understand or regulate corporate political influence through close-election-level candidate connections would have to start

somewhere other than the delivery channel our design captures. What investors, when they do update on close-election news, are pricing in legislator-identity information is itself a substantive question that the campaign-finance debate has not posed in these terms: to the extent any such price signal forms, the regulatory implication shifts from limiting firm-legislator delivery transactions to understanding how that signal forms and whom it advantages. A number of bounded caveats attach to this reading – short horizon, specific narrow channels below the resolution of our outcome data, legislator-level heterogeneity, the coarseness of the committee-to-issue-code crosswalk, and the possibility of strong-tie delivery in subpopulations below our detection threshold – and we report them at length in Appendix G. The paper’s central empirical contribution does not rest on resolving them. It is that in the setting where realized delivery should be most detectable, realized delivery does not show up in firm fundamentals or in firms’ own political behavior.

A Data Construction Details

A.1 PAC-GVKEY Linktable Extension to 2024

The Dane Christensen PAC-GVKEY link table (Christensen et al., 2022, 2023) covers cycles 1991–2020. We extend it through the 2022 and 2024 cycles by a Level-1 carry-forward that intersects the committee identifiers (CMTE_ID) active in the 2022 and 2024 Committee Master files with the 1991–2020 crosswalk, retaining rows whose CMTE_ID remained active in the new cycles. The extension increases the linktable from 54,102 rows to 60,406 rows. A substantive finding that emerges from the extension is that only 803 of the 1,577 firms in the 2000–2020 analysis panel (51%) retain an active corporate PAC in the 2022 or 2024 cycles – a substantial contraction that is not typically visible in finance-oriented political-connections studies. A Level-2 extension searching for new 2022/2024 CMTE_IDs not in the pre-2020 crosswalk identifies nine high-confidence matches, available as a merge-ready augmentation but not merged into the production linktable used

in the main regressions.

A.2 Panel Assembly for 2022 and 2024 Cycles

Panel extension to 2022–2024 adds 49,609 rows (23,410 in 2022, 26,199 in 2024) and 32 new firms. USASpending FPDS procurement coverage is pulled from local parquet to fiscal year 2026, with the FY 2026 window partial as of the analysis date. 2024 cycle rows carry `post_window_status = "truncated"` because the post-window $\{t+1, t+2\} = \{\text{FY 2025, FY 2026}\}$ is only partially observed. Main pooled regressions exclude the 2024 cycle.

A.3 Voteview Winner Enrichment

For pre-2022 cycles we attach the winning legislator’s `bioguide_id` and `icpsr` by joining MIT-Election-Lab race winners to the Voteview HSall roll-call member file (Congresses 107–118) on the key (state, district, Congress). The enrichment covers 99.6% of close-5pp general-election winners; the residual 0.4% failure set consists of cases where the MIT-EL and Voteview coding conventions disagree on at-large districts.

A.4 Committee-to-Issue-Code Crosswalk

Table 6 reports the hand-coded mapping from each of the eighteen standing committees of the U.S. House to the set of LDA issue codes associated with that committee’s jurisdiction. The mapping is one-to-many at the committee-to-issue-code level and is used in the lobbying target alignment analysis (§7.2).

Table 6: Committee-to-LDA-issue-code crosswalk.

Committee	LDA Issue Codes
Agriculture	AGR, FOO, ENV
Armed Services	DEF, AER, HOM, INT
Budget	BUD, FED
Education & Workforce	EDU, LBR, HCR
Energy & Commerce	ENG, ENV, HCR, CMN, CSP
Financial Services	FIN, BNK, INS, HOU
Foreign Affairs	FOR, INT, TRD
Homeland Security	HOM, INT, IMM
House Administration	GOV
Judiciary	JUS, CIV, IMM, IPR
Natural Resources	NAT, ENV, ENG, AGR
Oversight & Government Reform	GOV, BUD, FED
Rules	GOV
Science, Space, & Technology	SCI, CSP, TEC, ENG
Small Business	SMB, BNK, LBR
Transportation & Infrastructure	TRA, ENV, AGR, HOU
Veterans' Affairs	VET, HCR, LBR
Ways & Means	TAX, HCR, TRD, MED

Note: The crosswalk is hand-coded by the author from each committee's jurisdiction description in the House Rules. Issue codes follow LDA disclosure conventions.

B Composition Diagnostics

B.1 Paired-Permutation Balance

The paired-permutation balance test proceeds as follows: for each election, we draw one winner-supporting and one loser-supporting firm at random, compute the paired difference in the pre-window outcome, and average over the election sample; repeating the draw $K = 2,000$ times yields a permutation distribution of race-level paired balance that controls for within-race firm-selection noise. On the full pooled sample across 2000–2022, the test returns positive pre-period imbalance in 98.5% of draws for pre-procurement, with a mean paired difference of approximately +1.2 log-points. On the other eight outcomes the pattern is weaker but directionally similar: pre-window winner-loser paired differences are positive for every outcome except bond spreads, for which the sign reverses but the magnitude is small. The test is complementary to covariate-balance regressions in that it does not re-weight races by donor count and it does not collapse to the

race level.

B.2 Level-vs- Δ Decomposition

For each outcome we report the four-cell comparison of $\{\text{pre-level}, \text{post-level}, \Delta, \Delta^{\text{placebo}}\}$ coefficients on the win indicator. The pattern across outcomes is:

- **Procurement:** pre-level +0.30 (sig.), post-level +0.31 (sig.), Δ +0.006 (null), Δ^{placebo} +0.07 (null, pooled). Pure-level composition; differencing absorbs it cleanly, and the post-CU placebo is +0.107 ($p = 0.15$), also null.
- **ETR, bond spread, AAER violation:** all four cells null. No composition to correct.
- **Grants total assistance:** pre-level significantly positive, post-level null, Δ significantly negative, Δ^{placebo} significantly positive. Trajectory-reversal composition destroys identification.
- **Grants-only:** pre-level null, post-level significantly positive, Δ marginally positive, Δ^{placebo} borderline. Composition is small; Δ identifies, imperfectly.
- **Lobbying amount:** pre-level mildly positive (not significant), Δ null with clean placebo. Mild composition absorbed by differencing.
- **Target share:** pre-level and post-level both strongly positive at +0.10 ($p < 10^{-14}$), Δ null. Massive pre-existing composition absorbed by differencing.
- **PAC winner-party share:** placebo significantly positive, Δ significantly negative. Trajectory-reversal composition.

The decomposition pattern matches the Δ -RDD's theoretical design: differencing absorbs level composition cleanly; when trend bias is absent, Δ identifies; when trend bias is present (grants total assistance, grants any assistance, PAC winner-party share), Δ also fails and the placebo signals the failure.

B.3 Estimator Triangulation

Across the nine outcomes and three subperiods (pooled, pre-CU, post-CU), we report each Δ -RDD under clustered OLS, `rdrobust`, and Fisherian `rdrandinf`.

CAR₃ generals, baseline election-clustered specification. Clustered OLS $p = 0.005$, two-way (election $\times$ firm) cluster $p = 0.008$, Fisherian $p = 0.003$, collapsed-by-margin `rdrobust` $p = 0.13$: four inference frames agree on sign and rough magnitude of the +17.6 bps baseline point estimate. This within-specification four-frame agreement does not however extend across the Fowler et al. (2020) bandwidth $\times$ polynomial robustness grid on the same sample, under which only one of 32 specifications reaches conventional significance (Table 1; full 32-cell tables in Appendix C).

Three-frame agreement on Δ -RDD null outcomes. On procurement (pooled), ETR, bond spread, AAER violation, and lobbying amount, OLS, `rdrobust`, and `rdrandinf` all return null coefficients within each subperiod where the placebo is clean. Rare disagreements – procurement pre-CU OLS +0.14 versus `rdrobust` +0.10 ($p = 0.61$) and `rdrandinf` +0.30 ($p = 0.16$) – reflect mass-point pathology in the local-polynomial frame.

Three-frame disagreement on level procurement. Clustered OLS $p = 0.005$ on level-on-level procurement regression, Fisherian $p = 0.82$: the disagreement is by a factor of 160 and is the canonical signature of design failure.

Rare-binary degeneracy on the AAER indicators. `rdrobust` on AAER violation in the post-CU subperiod returns $t = -73.8$ with $SE = 0.0002$, a mass-point-pathology failure mode on a binary at the cutoff; the rarer AAER release indicator is degenerate in the same subperiod in the opposite direction, at $t = 9.3$ with $SE = 0.0005$. On neither indicator does the local-polynomial frame carry information post-CU, and on release the Fisherian frame disagrees with OLS on sign as well (Appendix D), so we read the post-CU AAER cells as cells the design does not identify rather than as estimates.

The triangulation rule is that we do not claim a causal effect when the three frames dis-

Table 7: FGS Table A8 grid on our panel, generals 2000–2020.

<i>Panel A: CAR₅ = CAR[−1, +5]; IK bandwidth = 0.00620 on half-margin scale</i>				
Polynomial	1pp	2.5pp	5pp	IK
Local constant (0)	−9.6 (29.7)	+27.8 (23.0)	+26.4 (17.3)	−1.5 (26.6)
Local linear (1)	+35.3 (51.0)	−19.7 (33.4)	+20.1 (28.2)	+0.9 (39.7)
Quadratic (2)	+35.8 (50.6)	−12.3 (42.4)	−28.8 (39.1)	+40.8 (59.7)
Cubic (3)	+35.8 (50.6)	−12.3 (42.4)	−43.4 (40.0)	+40.8 (59.7)
<i>N</i>	6,564	17,073	32,771	7,920
<i>Panel B: CAR₃ = CAR[−1, +1]; IK bandwidth = 0.00644 on half-margin scale</i>				
Polynomial	1pp	2.5pp	5pp	IK
Local constant (0)	−13.9 (24.5)	+23.1 (18.2)	+29.0* (13.7)	−4.9 (23.6)
Local linear (1)	−23.9 (42.2)	−27.4 (28.1)	+9.0 (23.4)	−33.4 (31.1)
Quadratic (2)	−23.9 (42.0)	−51.3 (34.4)	−34.9 (31.1)	−22.5 (50.8)
Cubic (3)	−23.9 (42.0)	−51.3 (34.4)	−42.7 (32.1)	−22.5 (50.8)
<i>N</i>	6,564	17,073	32,771	8,156

Note: Point estimates in basis points with firm-clustered standard errors in parentheses; * $p < .05$, ** $p < .01$. Grid counts: Panel A = 7 of 16 negative, 0 sig. positive, 0 sig. negative. Panel B = 13 of 16 negative, 1 sig. positive (local constant $\times$ 5pp), 0 sig. negative. Combined = 20 of 32 negative, 1 of 32 sig. positive, 0 of 32 sig. negative. Double-donor firms dropped (FGS convention).

agree. A frame-disagreement null is not read as a causal zero either; it is read as “design does not identify here.”

C Event-Window CAR Robustness Grid

Table 7 reports the full 32-cell Fowler et al. (2020) Table A8 grid on our panel ($N = 32,771$ firm-elections, 1,237 firms, 554 elections, generals 2000–2020). Table 8 reports the same grid on Akey’s bundled data ($N = 235$ firm-elections, 95 firms). Both samples drop double-donor firms, use firm-clustered standard errors, and use no fixed effects, faithful to the FGS Stata specification (`qpq_Appendix_Tables.do` lines 1626–1858). IK bandwidth on our panel is computed via CCT-MSERD (`rdrobust`), yielding ≈ 0.006 on the half-margin scale; on Akey’s data, FGS’s hardcoded IK values (0.017 for Panel A, 0.0119 for Panel B) are used. Point estimates are in basis points; standard errors in parentheses.

Table 8: FGS Table A8 grid on Akey (2015) bundled data (benchmark).

<i>Panel A: CAR₅ = CAR[−1, +5]; IK bandwidth = 0.01700 (FGS-hardcoded)</i>				
Polynomial	1pp	2.5pp	5pp	IK
Local constant (0)	+169 (154)	+134 (91)	+5 (64)	+64 (71)
Local linear (1)	+753 (632)	+463 (233)	+300* (122)	+285* (139)
Quadratic (2)	+753 (632)	+458 (255)	+683** (246)	+425 (298)
Cubic (3)	+753 (632)	+458 (255)	+683** (246)	+425 (298)
<i>N</i>	34	77	234	124
<i>Panel B: CAR₃ = CAR[−1, +1]; IK bandwidth = 0.01190 (FGS-hardcoded)</i>				
Polynomial	1pp	2.5pp	5pp	IK
Local constant (0)	−39 (91)	−56 (61)	−41 (40)	−56 (61)
Local linear (1)	+244 (326)	+203 (138)	−15 (87)	+203 (138)
Quadratic (2)	+244 (326)	+190 (145)	+369* (150)	+190 (145)
Cubic (3)	+244 (326)	+190 (145)	+369* (150)	+190 (145)
<i>N</i>	34	77	234	77

Note: Point estimates in basis points with firm-clustered standard errors in parentheses; * $p < .05$, ** $p < .01$. Grid counts: 5 of 32 negative point estimates, 6 sig. positive, 0 sig. negative. Akey’s published +3% headline corresponds to the Panel A $\times$ local linear $\times$ 5pp cell (+300 bps, $p = 0.016$). Double-donor firms dropped (FGS convention).

D Other Government-Benefit Channels – Detail

D.1 Effective Tax Rate

GAAP effective tax rate is $\text{etr_gaap} = \text{txt}/\text{pi}$, winsorized to $[0, 1]$. The pooled Δ -RDD returns +0.0012 ($p = 0.63$); the placebo +0.0013 ($p = 0.61$). The 95% confidence interval of approximately $[-0.004, +0.006]$ rules out ETR responses larger than about half a percentage point. Running the same regression in levels returns -0.0006 ($\text{SE} = 0.0046$; $p = 0.90$) on the post window and -0.0005 ($\text{SE} = 0.0043$; $p = 0.91$) on the pre window: neither level is imbalanced. ETR is not a size-correlated outcome; for this outcome Δ -RDD is a precision-improving re-expression rather than a composition-correcting device. Pre-CU (-0.005 , $p = 0.13$) and post-CU ($+0.005$, $p = 0.14$) are separately null under OLS. A small pre-CU `rdrobust` placebo signal is not echoed by OLS or Fisherian inference, consistent with local-polynomial mass-point pathology. Substantive reading: effective tax rates are dominated by statutory rates, the IRS audit regime, and corporate tax-planning

infrastructure, not by the identity of a marginally elected House member within a two-year window.

D.2 Bond Spread

The outcome is `spread_wmean`, the issuance-amount-weighted within-firm-year mean corporate bond spread to comparable-duration Treasury, winsorized to $[0, 0.05]$. Coverage is restricted to the 61% of main-sample firms that are public-debt issuers; the sample contains 24,586 close-5pp firm-elections from 541 elections. The pooled Δ -RDD returns -0.00024 ($p = 0.25$) with a clean placebo. The point estimate corresponds to roughly 2.4 annualized basis points of compression, well below the 20–50 basis-point compression magnitudes documented in the prior firm-connection credit-pricing literature and not statistically distinguishable from zero. Level-RDD specifications on pre and post are individually null; no composition correction is operative. The sign-consistency with the compression hypothesis, paired with the inability to distinguish the magnitude from zero, is consistent with the null-persistence pattern on the equity side (§5): the bond market, like the equity market beyond the three-day event window, does not price a sustained connection premium.

D.3 SEC Enforcement (AAER)

The primary AAER outcome is $\Delta I(\text{any event})$, taking values in $\{-1, 0, +1\}$. AAER is rare: 155 of 1,577 main-sample firms are ever flagged, and the within-window any-event rate is 2–3%. Pooled: -0.0010 ($p = 0.62$); placebo $+0.0001$ ($p = 0.95$). 95% confidence interval approximately $[-0.005, +0.003]$. Pre-CU (-0.0015 , $p = 0.70$) and post-CU (-0.0010 , $p = 0.63$) are separately null under OLS. A release-year sensitivity indicator (four times rarer than the violation-period cells) is likewise null under OLS in every subperiod, post-CU included at -0.0019 ($p = 0.15$) against a clean placebo ($p = 0.35$). The other two frames

Table 9: Crisis $\times$ win interaction on four fundamentals outcomes.

Outcome	Tight crisis (GFC + COVID)			Broad crisis (+ dotcom, 9/11)		
	β_3	p (main)	p (plac.)	β_3	p (main)	p (plac.)
Procurement	−0.087	0.29	0.49	+0.087	0.26	0.62
ETR	+0.005	0.34	0.036	−0.000	0.98	0.005
Bond spread	+0.00040	0.21	0.23	−0.00025	0.51	0.23
AAER violation	+0.0036	0.42	0.70	+0.0035	0.39	0.20
AAER release	−0.0012	0.56	0.046	+0.0009	0.67	0.023

do not agree with OLS on that post-CU cell: Fisherian `rdrandinf` returns +0.0180 with $p = 0.008$ – significant, and of the opposite sign – and `rdrobust` returns $t = 9.3$ with $SE = 0.0005$, a rare-binary degeneracy. Under our triangulation rule, three frames that disagree on sign mark a cell the design does not identify; they do not evidence a leniency effect. We flag conditional-logit or matched-event-study specifications as deferred follow-up.

E Off-Equilibrium Preemption – Detail

E.1 Crisis $\times$ Win Interaction

We estimate

$$\Delta y_i = \alpha + \beta_1 \cdot \text{win}_i + \beta_2 \cdot \text{crisis_post}_i + \beta_3 \cdot (\text{win}_i \times \text{crisis_post}_i) + \varepsilon_i, \quad (4)$$

on four fundamentals outcomes plus the AAER release variant. The tight crisis definition (GFC peak + COVID) assigns $\text{crisis_post} = 1$ to cycles $\{2006, 2008, 2018\}$; the broad definition adds $\{2000, 2002\}$. For each outcome and each definition we report both the main interaction and the placebo interaction. Table 9 reports the result.

No main interaction reaches conventional significance; p -values range from 0.21 to 0.98. Four placebo interactions, however, fail at conventional levels: ETR under both cri-

sis definitions (tight $p = 0.036$, broad $p = 0.005$) and AAER release under both (tight $p = 0.046$, broad $p = 0.023$). On those two outcomes the between-winner and between-loser composition is itself state-dependent, so a main-interaction estimate would be indistinguishable from composition artifact even if it were statistically significant. On procurement, bond spread, and AAER violation the placebo interactions are clean, so the null readings on the main interactions are direct evidence of no crisis-state delivery.

E.2 Industry-Targeted Scenarios

Industry-targeted interactions cross three scenarios (Defense $\times$ 2002; Finance $\times$ 2006–2008; Finance $\times$ 2018) with the same five outcomes used above, producing fifteen main interaction tests and fourteen placebo tests. Two main coefficients reach $p < 0.05$: Defense $\times$ 2002 $\times$ ETR (-0.074 , $p = 0.036$) and Defense $\times$ 2002 $\times$ AAER violation (-0.041 , $p = 0.032$). Fifteen uncorrected tests at the five-percent level produce 0.75 false positives in expectation, so a count of two significant cells is inside the range a chance pattern generates. Neither placebo coefficient is significant, and both cells are thin: the ETR interaction is estimated on 2,881 firm-elections of which 235 fall in the crisis cell, the AAER-violation interaction on 3,247 of which 240 fall in the crisis cell, across 36 close elections in each case. We document the signals as suggestive rather than as findings. What this appendix rules out is a broad-stroke insurance interpretation on the outcomes we measure. What it does not rule out is insurance operating through specific narrow channels – TARP Capital Purchase Program allocations, no-bid COVID-era emergency procurement captured at the sub-program level within FPDS, or regulatory waivers and emergency use authorizations not aggregated into the AAER feed.

F Supplementary Tables and Additional Robustness

F.1 Hedging Rate Over Time

The corporate PAC hedging rate – share of firm-election observations in close-5pp races where the firm contributed to both candidates – fell from 4.3% in the 2010 cycle to 0.6% in the 2022 cycle. The decline is progressive rather than discrete: 4.3% (2010), 4.0% (2012), 2.1% (2014), 1.2% (2016), 0.79% (2018), 0.67% (2020), 0.6% (2022). The inflection is around 2014–2018 rather than at Citizens United itself.

F.2 Within-Firm Fixed Effects

An alternative to Δ -RDD for absorbing firm-specific composition is within-firm fixed effects, identifying off the 1,047 firms that appear on both winner and loser sides across different close races:

$$\log(1 + \text{procurement}_{i,\text{post}}) = \alpha_i + \delta_t + \gamma \cdot \log(1 + \text{procurement}_{i,\text{pre}}) + \beta \cdot \text{win}_i + \varepsilon_i. \quad (5)$$

Estimated on 34,375 dual-appearance firm-elections from 588 close elections, the pooled coefficient is $\hat{\beta} = -0.002$ with $\text{SE} = 0.025$ and $p = 0.92$, a 95% confidence interval of $[-0.052, +0.047]$. The subperiods are separately null: pre-CU $+0.048$ ($p = 0.15$, interval $[-0.018, +0.113]$) and post-CU -0.044 ($p = 0.16$, interval $[-0.105, +0.018]$). Every subperiod interval excludes procurement gains at the magnitudes the adjacent literature documents. The within-firm FE specification is independent of the Δ -RDD in that it uses only the dual-appearance subpopulation and absorbs firm-specific level via fixed effects rather than differencing. Convergence to a tightly estimated zero is a useful robustness check.

F.3 Special-Election Augmentation

We assembled a Wikipedia-scraped special-election panel covering all House special elections from 2002 to 2024, matched candidates to FEC candidate master, and attached PAC-GVKEY linkage and USASpending procurement. The augmented close-5pp panel contains 67 close special elections in addition to the general-election sample. On the primary procurement specification the general-only estimate is +0.153; the augmented general-plus-special estimate is +0.112. The augmentation attenuates rather than amplifies, and the $\text{win} \times \text{is_special}$ interaction under clustered SE returns -0.82 with $p = 0.13$ – directionally opposite but underpowered. A pipeline-specific-construction-error interpretation would predict amplification. The attenuation pattern is consistent with a general across-subtype null.

F.4 Connection-Strength Robustness

We rerun the main Δ -RDD procurement regression on the firms whose connection to their candidate is strongest. Connection strength is the firm’s cumulative pre-window PAC contribution to the candidate defining its side of the close race – the candidate it backed, whether that candidate went on to win or lose. Measuring the proxy on the firm’s own candidate rather than on the race winner keeps it pre-treatment and symmetric across the two sides, which conditioning on money given to the eventual winner would not be. The top decile is a contribution of \$5,000 or more; it retains 6,446 firm-elections from 667 firms and 549 close elections, 4,165 of them winner-side and 2,281 loser-side. On that subpopulation the pooled Δ -RDD returns +0.022 (SE = 0.086; $p = 0.80$) against a clean pre-trend placebo (+0.146, $p = 0.25$), with a 95% confidence interval of $[-0.15, +0.19]$. The point estimate is within sampling error of the full-sample +0.006; the interval widens with the reduced sample but still excludes procurement gains at the magnitudes the adjacent literature documents. Subperiod estimates are separately null: pre-CU +0.215 ($p = 0.16$) and

post-CU -0.070 ($p = 0.50$). No evidence of strong-tie delivery emerges in this subpopulation.

F.5 Portfolio Breadth Stratification

A potential reader critique of the Δ -RDD null is that firm-level outcomes aggregate across a diversified political portfolio, so a single close-race win/loss is a small share of a firm's total political exposure and the firm-level RDD effect is mechanically attenuated under portfolio diversification regardless of whether connections deliver per race. If this cross-race dilution mechanism is operative, two predictions follow. First, restricting to single-candidate firm-cycles – the maximally concentrated slice where no dilution can occur – should recover a positive delivery effect. Second, the $\text{win} \times \log(n_{\text{cands}})$ interaction should be negative and significant, since the per-firm treatment effect should shrink monotonically with portfolio breadth. We test both predictions across all nine primary outcomes.

Firm-cycle portfolio breadth is the number of distinct House candidates a firm contributed to in a given calendar year. In our combined 2000–2024 panel the median firm-cycle contains 10 candidates (mean 21.8); only 10.1% of firm-cycles involve a single candidate, and single-candidate firm-cycles account for under 0.2% of total PAC dollars across the broader FEC universe of publicly traded firms. The modal firm is diversified across many races; single-candidate firms are atypical and small.

Across all nine outcomes the dilution prediction fails. Six interaction coefficients are null with p -values above 0.20 (procurement -0.006 , $p = 0.84$; AAER violation -0.002 , $p = 0.22$; AAER release -0.001 , $p = 0.52$; log lobbying \$ -0.015 , $p = 0.51$; any lobbying -0.001 , $p = 0.55$; grants-only -0.16 , $p = 0.22$). Three are signed opposite to the dilution prediction (ETR $+0.0008$, $p = 0.62$; bond spread $+0.0001$, $p = 0.14$; grants total $+0.10$, $p = 0.53$). Not one outcome exhibits the negative-significant interaction the dilution hypothesis requires.

Single-candidate subsample estimates are null or negative for every outcome. The only nominally positive point estimate is AAER violation ($+0.019$, $p = 0.088$, $N = 204$),

which points toward *more* enforcement for connected firms – the opposite direction from any “connections reduce enforcement” reading. Procurement single-candidate: -0.330 , SE 0.409 , $p = 0.42$, $N = 204$; 95% confidence interval approximately $[-1.1, +0.5]$. The interval rules out large positive delivery effects at the magnitudes some CAR-literature readings imply, though it is too wide to rule out small effects at the magnitude of the equity event-window reaction itself. Grants total shows non-monotonic breadth heterogeneity – bin 4–10: -2.03 , $p = 0.008$; bin 31+: -1.07 , $p < 0.001$; bin 2–3: $+3.50$, $p = 0.07$ – inconsistent with a monotonic dilution mechanism and more plausibly reflecting a measure-scale artifact from combining grants, loans, and insurance; the grants-only specification returns null across all five breadth bins.

The delivery null is robust to portfolio breadth across every outcome in the paper. Firms’ diversified portfolio structure is a descriptive feature of how corporate political investment is organized, not a statistical confound that rescues the null interpretation.

F.6 One-Close-Race Firm-Cycles

The cross-race aggregation concern of §4.4 is that a firm-cycle touching several close races carries a mixture of wins and losses into a single Δy_i . This subsection re-estimates the Δ -RDD on the firm-cycles that contain exactly one close race, where no mixture is possible by construction. The subset rule counts close races ($|\text{margin}| \leq 0.05$) per firm and calendar year on the analysis sample – close-5pp, hedgers dropped, 2024 excluded – and keeps the firm-cycles whose count is one; applying the same rule before the hedger drop returns 1,797 rather than 1,886 firm-elections. The slice retains 1,886 firm-elections from 986 firms and 468 elections, out of 5,795 close-5pp firm-cycles: roughly nine times the single-candidate slice of Appendix F.5, because a firm may fund many races of which only one is close. Table 10 reports all nine primary outcomes on the slice.

Seven of the nine outcomes are null with clean placebos, procurement among them: -0.174 ($p = 0.19$) with placebo $p = 0.88$ and a 95% confidence interval of $[-0.44, +0.09]$,

Table 10: Δ -RDD on one-close-race firm-cycles, pooled 2000–2022.

Outcome	$\hat{\beta}$	SE	p -value	Placebo p	N
Procurement ($\Delta \log(1 + \$)$)	−0.174	0.134	0.19	0.88	1,830
ETR (Δ rate)	−0.011	0.009	0.23	0.57	1,447
Bond spread (Δ rate)	−0.00040	0.00039	0.31	0.33	870
AAER violation ($\Delta I(\text{any})$)	−0.0075	0.0047	0.12	0.007	1,882
Grants only (excl. loans)	+0.079	0.090	0.38	0.72	1,886
Lobbying $\Delta \log(1 + \$)$	−0.052	0.202	0.80	0.48	1,886
Lobbying $\Delta \log(1 + n_{\text{reports}})$	−0.026	0.034	0.46	0.64	1,886
Target count $\Delta \log(1 + \cdot)$	+0.006	0.023	0.80	0.47	1,686
PAC Δ winner-party share	−0.061	0.031	0.046	0.22	1,142

Note: Clustered OLS on the win indicator, election-clustered standard errors, on firm-cycles containing exactly one close race ($|\text{margin}| \leq 0.05$, hedgers dropped, 2024 cycle excluded). Windows follow the conventions of §3. Per-outcome N varies with outcome coverage: the target-count outcome requires the committee crosswalk (425 elections on this slice), and the PAC outcome uses cycle-based windows ending with the 2020 cycle.

against $[-1.1, +0.5]$ on the single-candidate slice. The procurement subperiods are separately null with clean placebos (pre-CU -0.048 , $p = 0.81$; post-CU -0.265 , $p = 0.13$). The two remaining cells require qualification. AAER violation’s placebo fails on this slice ($p = 0.007$), so under the placebo rule that cell does not identify here, and we draw no conclusion from its null point estimate. PAC winner-party share is -0.061 ($p = 0.046$) with a clean placebo ($p = 0.22$): the slice recovers, with clean identification, the winner-party-share decline that fails its placebo on the full battery (§7.3). It is one significant cell in nine tests, but its content matters for interpretation: it is a political-behavior outcome signed toward diversification, not a benefit outcome signed toward delivery. Restricting to mixture-free firm-cycles therefore leaves the delivery null intact on every government-benefit outcome and, on the one cell that moves, supports the portfolio reading of firm behavior.

F.7 Alternative Sample Restrictions

Relaxing the hedger exclusion to include all firms in close races (assigning each to the winner or loser side by the sign of its net donation) does not change the direction or

significance of any of the primary outcomes. A uniform 32-cell A-versus-B sensitivity test records each headline analysis under both the drop-hedgers and the retain-hedgers-assign-by-side conventions: eleven CAR event, drift, and placebo windows; the CAR event window on the pre-Citizens-United and post-Citizens-United subsamples; pooled, pre-CU, and post-CU bond-spread Δ -RDDs with a placebo; pooled, pre-CU, and post-CU procurement Δ -RDDs; and ETR, AAER, grants, and lobbying Δ -RDDs with placebos. All 32 of 32 analyses pass a predefined materiality criterion (same sign, point-estimate change $\leq 30\%$, p -value does not cross 0.05; small-null exemption for placebos where $|t| < 1$ and $p > 0.3$ under both conventions). The event-window CAR estimate specifically shifts by -1% between conventions ($+17.58$ bps under drop-hedgers versus $+17.41$ bps under retain-by-side, both at $p \approx 0.005$); the gap between the drop-hedgers baseline and the $+29$ bps FGS-convention estimate in §5 is therefore not driven by hedger handling but by kernel weighting and standard-error clustering. Tightening the bandwidth to $|\text{margin}| \leq 0.03$ reduces precision across all outcomes but does not change the sign of any point estimate. Including the 2024 cycle yields coefficients within sampling error of the 2000–2022 pooled estimates.

F.8 Full Triangulation Table

For each of the nine primary outcomes we report coefficients, standard errors, p -values, and effective N under clustered OLS, `rdrobust` with triangular kernel and mass-point adjustment, and Fisherian `rdrandinf` at $|\text{margin}| \leq 0.05$, on the pooled 2000–2022 sample, the pre-CU 2000–2008 subsample, and the post-CU 2010–2022 subsample. Table 11 reports the full nine-outcome by three-subperiod by three-estimator grid, generated directly from the canonical results table by `code/build_triangulation_table.R`.

Table 11: Full triangulation grid: nine primary outcomes, three subperiods, three estimators.

Subperiod	Clustered OLS			rdrobust			rdrandinf	
	N	$\hat{\beta}$ (SE)	p	$\hat{\beta}$ (SE)	p	N_{eff}	$\hat{\beta}$	p
<i>Procurement ($\Delta \log(1 + \\$)$)</i>								
Pooled 2000–2022	35,033	+0.0065 (0.045)	0.885	−0.179 (0.374)	0.502	8,265	+0.049	0.714
Pre-CU 2000–2008	14,036	+0.139 (0.064)	0.030	+0.096 (0.256)	0.610	2,964	+0.299	0.161
Post-CU 2010–2022	20,997	−0.089 (0.055)	0.103	−0.325 (0.562)	0.414	6,552	−0.117	0.510
<i>Effective tax rate (Δ rate)</i>								
Pooled 2000–2022	31,703	+0.0012 (0.0024)	0.632	+0.010 (0.0094)	0.201	8,944	+0.0015	0.815
Pre-CU 2000–2008	12,545	−0.0047 (0.0031)	0.128	−0.0058 (0.015)	0.874	2,902	+0.013	0.174
Post-CU 2010–2022	19,158	+0.0050 (0.0034)	0.143	+0.025 (0.016)	0.061	7,672	−0.0069	0.475
<i>Bond spread (Δ rate)</i>								
Pooled 2000–2022	24,586	−0.00024 (0.00021)	0.250	+0.00042 (0.00086)	0.849	6,521	+0.00033	0.107
Pre-CU 2000–2008	8,319	−0.00034 (0.00049)	0.482	+0.00001 (0.00081)	0.730	1,523	+0.00013	0.769
Post-CU 2010–2022	16,267	−0.00017 (0.00016)	0.288	+0.00059 (0.00096)	0.378	7,172	+0.00049	0.014
<i>AAER violation ($\Delta I(\text{any})$)</i>								
Pooled 2000–2022	35,611	−0.0010 (0.0021)	0.623	+0.0044 (0.029)	0.982	10,611	+0.0036	0.693
Pre-CU 2000–2008	14,392	−0.0015 (0.0038)	0.697	+0.011 (0.051)	0.848	6,519	+0.017	0.131
Post-CU 2010–2022	21,219	−0.0010 (0.0021)	0.627	−0.015 (0.00024)	< 0.001	1,810	−0.0062	0.579
<i>Grants only (excl. loans) ($\Delta \log(1 + \\$)$)</i>								
Pooled 2000–2022	35,623	+0.053 (0.026)	0.041	−0.353 (0.419)	0.334	8,495	−0.043	0.677
Pre-CU 2000–2008	14,392	+0.103 (0.048)	0.031	−2.402 (0.527)	< 0.001	2,410	−0.057	0.755
Post-CU 2010–2022	21,231	+0.015 (0.027)	0.568	+0.273 (0.366)	0.327	9,509	−0.023	0.825
<i>Lobbying expenditure ($\Delta \log(1 + \\$)$)</i>								
Pooled 2000–2022	35,623	−0.093 (0.064)	0.144	+0.361 (0.113)	< 0.001	5,600	−0.060	0.530
Pre-CU 2000–2008	14,392	+0.020 (0.030)	0.501	+0.184 (0.035)	< 0.001	2,111	+0.173	0.166
Post-CU 2010–2022	21,231	−0.178 (0.102)	0.080	+0.246 (0.053)	< 0.001	1,257	−0.216	0.086
<i>Lobbying reports ($\Delta \log(1 + n)$)</i>								
Pooled 2000–2022	35,623	−0.013 (0.011)	0.242	+0.101 (0.017)	< 0.001	7,184	+0.0024	0.893
Pre-CU 2000–2008	14,392	−0.016 (0.015)	0.308	+0.094 (0.043)	0.052	3,937	+0.040	0.097
Post-CU 2010–2022	21,231	−0.014 (0.0099)	0.167	+0.086 (0.033)	0.005	2,502	−0.020	0.322
<i>Target-domain reports ($\Delta \log(1 + n)$)</i>								
Pooled 2000–2022	33,507	−0.044 (0.013)	< 0.001	+0.055 (0.061)	0.364	13,349	−0.027	0.311
Pre-CU 2000–2008	13,815	−0.0090 (0.018)	0.611	+0.161 (0.090)	0.100	5,637	+0.054	0.138
Post-CU 2010–2022	19,692	−0.072 (0.014)	< 0.001	−0.020 (0.074)	0.770	7,693	−0.083	0.017
<i>PAC winner-party share (Δ share)</i>								
Pooled 2000–2022	26,010	−0.097 (0.018)	< 0.001	−0.136 (0.041)	< 0.001	6,465	−0.061	< 0.001
Pre-CU 2000–2008	9,267	−0.155 (0.029)	< 0.001	−0.152 (0.064)	0.028	1,972	−0.090	< 0.001
Post-CU 2010–2022	16,743	−0.062 (0.021)	0.004	−0.118 (0.056)	0.014	4,271	−0.046	< 0.001

Note: Main (non-placebo) Δ -RDD coefficients on the win indicator at $|\text{margin}| \leq 0.05$, hedging firms excluded, 2024 cycle excluded. Clustered OLS uses election-clustered standard errors on the full subperiod sample; *rdrobust* is the local-polynomial estimator of Calonico et al. (2014) with triangular kernel and mass-point adjustment, reported with its effective sample N_{eff} ; *rdrandinf* is the Fisherian local-randomization estimator of Cattaneo et al. (2016), which returns a randomization p -value and no standard error. Two outcomes end earlier than the stated windows by construction and their rows are labelled with the common window for comparability: target-domain reports require the winner committee crosswalk, which covers Congresses 107–117, and the PAC outcomes use cycle-based windows that need post cycle $t + 2$ to be observed, so both end with the 2020 cycle. Cells are generated directly from `results/RESULTS_TABLE.csv` by `code/build_triangulation_table.R`.

G Scope Conditions

The null results in §6 and §7 are bounded in several ways that readers should keep in mind.

Horizon. Our post-election window is two fiscal years, aligned with typical committee-assignment stability and with standard event-study practice. Effects that operate on longer horizons – goodwill accumulated over a legislator’s full tenure, slow-building regulatory posture shifts that materialize only after several Congresses, long-run credit-reputation advantages – are not within our detection range. The cumulative-abnormal-return literature’s implicit pricing horizon is also short-run, so a short-run null on realized delivery is an appropriate test of the reaction’s implicit content; a longer-horizon delivery story is a distinct empirical claim that would need a different design.

Specific narrow channels. Our crisis-state tests operationalize the insurance alternative at the election-cycle level rather than at the specific-program level. TARP Capital Purchase Program allocations in 2008–2009, no-bid COVID-era emergency procurement captured at the sub-program level within FPDS, and regulatory waivers or emergency-use authorizations not aggregated into AAER are all potential channels through which a state-contingent connection-delivery mechanism could operate. Our design cannot detect these channels; a focused study on each specific program would be needed.

Legislator-level heterogeneity. Our design aggregates across all close-election winners regardless of the winner’s committee assignment, seniority, or party control. A legislator-heterogeneity analysis – winners on the House Financial Services or Armed Services committees, or winners in majority-party seats versus minority-party seats – might uncover pockets of delivery that the aggregate masks. We view this as a natural follow-up rather than as a limitation of the present paper’s central claim about average delivery.

Lobbying-target granularity. The committee-to-issue-code crosswalk used for target alignment in §7.2 is hand-coded at the standing-committee level and does not reflect sub-committee jurisdictions or individual-bill sponsorship. A bill-level targeting measure –

firms lobbying on specific bills sponsored or co-sponsored by the connected legislator – would be sharper.

LDA historical coverage. LDA filings begin in 1999, and the intensive-margin lobbying data for cycles 2000 and 2002 is thin. Pre-CU lobbying analyses use four cycles, three of which overlap the LDA ramp-up.

Connection strength. We define a firm’s connection to a candidate through any corporate PAC contribution above the FEC disclosure threshold, following Akey (2015)’s close-election design. A richer measure – size of contribution, multi-cycle persistence of the donating relationship, board-level personal ties – might identify a subpopulation of strong-tie firm-legislator dyads where delivery does occur. The top-decile-connection-strength specification in Appendix F returns comparable nulls with wider confidence intervals. The paper does not rule out delivery for a narrow strong-tie subpopulation it does not observe; it does rule out delivery operating at the average close-election donation that the literature has treated as the identifying unit.

References

- Agca, S. and Igan, D. (2023). The Lion's Share: Evidence from Federal Contracts on the Value of Political Connections. *Journal of Law and Economics*, 66(3). Semantic Scholar ID: e2846de11f194c3bed5ff6e9d25d33dcf5cca893.
- Akey, P. (2015). Valuing Changes in Political Networks: Evidence from Campaign Contributions to Close Congressional Elections. *Review of Financial Studies*, 28(11):3188–3223.
- Ansolabehere, S., de Figueiredo, J. M., and Snyder Jr., J. M. (2003). Why Is There So Little Money in U.S. Politics? *Journal of Economic Perspectives*, 17(1):105–130.
- Baltrunaite, A. (2020). Political contributions and public procurement: Evidence from Lithuania. *Journal of the European Economic Association*, 18(2):541–582.
- Barber, B. M. and Lyon, J. D. (1997). Detecting Long-run Abnormal Stock Returns: The Empirical Power and Specification of Test Statistics. *Journal of Financial Economics*, 43(3):341–372.
- Barber, M. J. (2016). Donation Motivations: Testing Theories of Access and Ideology. *Political Research Quarterly*, 69(1):148–159.
- Bertoli, A. and Hazlett, C. (2025). Seeing Like a District: Understanding What Close-Election Designs for Leader Characteristics Can and Cannot Tell Us. *Political Analysis*, 33(4):393–411.
- Bertrand, M., Bombardini, M., and Trebbi, F. (2014). Is It Whom You Know or What You Know? An Empirical Assessment of the Lobbying Process. *American Economic Review*, 104(12):3885–3920.
- Bisetti, E., Lewellen, S., Sarkar, A., and Zhao, X. (2026). Smokestacks and the Swamp. *Review of Financial Studies*, 39(9).

- Bonica, A. (2016). Avenues of Influence: On the Political Expenditures of Corporations and Their Directors and Executives. *Business and Politics*, 18(4):367–394.
- Boubakri, N., Cosset, J.-C., and Saffar, W. (2012). The Impact of Political Connections on Firms’ Operating Performance and Financing Decisions. *Journal of Financial Research*, 35(3):397–423.
- Calonico, S., Cattaneo, M. D., and Titiunik, R. (2014). Robust nonparametric confidence intervals for Regression-Discontinuity Designs. *Econometrica*, 82(6):2295–2326.
- Cattaneo, M. D., Idrobo, N., and Titiunik, R. (2024). *A Practical Introduction to Regression Discontinuity Designs: Extensions*. Cambridge Elements in Quantitative and Computational Methods for the Social Sciences. Cambridge University Press.
- Cattaneo, M. D., Titiunik, R., and Vazquez-Bare, G. (2016). Inference in Regression Discontinuity Designs under Local Randomization. *Stata Journal*, 16(2):331–367.
- Cattaneo, M. D., Titiunik, R., and Vazquez-Bare, G. (2020). Analysis of Regression-Discontinuity Designs with Multiple Cutoffs or Multiple Scores. *Stata Journal*, 20(4):866–891.
- Christensen, D. M., Jin, H., Sridharan, S., and Wellman, L. A. (2022). Hedging on the Hill: Does political hedging reduce firm risk? *Management Science*, 68(6):3975–4753. Maintains the PAC-to-GVKEY link table (*PAC_GVKEY_Linktable_1991_2020*) released at <https://sites.google.com/view/danechristensen/data/fec-to-compustat-link-table>.
- Christensen, D. M., Morris, A. S., Walther, B. R., and Wellman, L. A. (2023). Political information flow and management guidance. *Review of Accounting Studies*, 28(3):1466–1499.

- Cohen, L., Diether, K., and Malloy, C. (2013). Legislating Stock Prices. *Journal of Financial Economics*, 110(3):574–595.
- Correia, M. M. (2014). Political Connections and SEC Enforcement. *Journal of Accounting and Economics*, 57(2–3):241–262.
- De Magalhães, L., Hangartner, D., Hirvonen, S., Meriläinen, J., Ruiz, N. A., and Tukiainen, J. (2025). When Can We Trust Regression Discontinuity Design Estimates from Close Elections? Evidence from Experimental Benchmarks. *Political Analysis*, 33(3):258–265.
- Denes, M. and Scanlon, M. M. (2025). Shining a Light on Firms’ Political Connections: The Role of Dark Money. *Review of Corporate Finance Studies*, 14(4):989–1023.
- Duchin, R. and Sosyura, D. (2012). The Politics of Government Investment. *Journal of Financial Economics*, 106(1):24–48.
- Eggers, A. C. and Hainmueller, J. (2013). Capitol Losses: The Mediocre Performance of Congressional Stock Portfolios. *Journal of Politics*, 75(2):535–551.
- Faccio, M. (2006). Politically Connected Firms. *American Economic Review*, 96(1):369–386.
- Fazekas, M., Ferrali, R., and Wachs, J. (2022). Agency independence, campaign contributions, and Favouritism in US Federal Government Contracting. *Journal of Public Administration Research and Theory*, 32(4).
- Fotak, V., Lee, H. S., Megginson, W. L., and Salas, J. M. (2025). The Political Economy of Tariff Exemption Grants. *Journal of Financial and Quantitative Analysis*, 60(6):2678–2717.
- Fournaies, A. and Fowler, A. (2021). Do Campaign Contributions Buy Favorable Policies? Evidence from the Insurance Industry. *Political Science Research and Methods*.
- Fournaies, A. and Hall, A. B. (2014). The Financial Incumbency Advantage: Causes and Consequences. *Journal of Politics*, 76(3):711–724.

- Fourinaies, A. and Hall, A. B. (2018). How Do Interest Groups Seek Access to Committees? *American Journal of Political Science*, 62(1):132–147.
- Fowler, A., Garro, H., and Spenkuch, J. L. (2020). Quid Pro Quo? Corporate Returns to Campaign Contributions. *Journal of Politics*, 82(3):844–858.
- Goldman, E., Rocholl, J., and So, J. (2009). Do Politically Connected Boards Affect Firm Value? *Review of Financial Studies*, 22(6):2331–2360.
- Goldman, E., Rocholl, J., and So, J. (2013). Political connections and the Allocation of Procurement Contracts. *Review of Finance*, 17(5):1617–1648.
- Grier, K. B., Munger, M. C., and Roberts, B. E. (1994). The Determinants of Industry Political Activity, 1978–1986. *American Political Science Review*, 88(4):911–926.
- Hall, R. L. and Wayman, F. W. (1990). Buying Time: Moneyed Interests and the Mobilization of Bias in Congressional Committees. *American Political Science Review*, 84(3):797–820.
- Hartman, E. (2021). Equivalence Testing for Regression Discontinuity Designs. *Political Analysis*, 29(4):505–521.
- Houston, R., Ferris, S. P., and Pérez-Amuedo, J. A. (2026). All Politics Is Local: Corporate Political Power and the Award of Federal Contracts. *Contemporary Economic Policy*. Early View.
- Jayachandran, S. (2006). The Jeffords Effect. *Journal of Law and Economics*, 49(2):397–425.
- Kim, I. S., Stuckatz, J., and Wolters, L. (2026). Systemic and Sequential Links Between Campaign Donations and Lobbying. *Journal of Politics*, 88(1):110–124.
- Kothari, S. P. and Warner, J. B. (2007). Econometrics of Event Studies. In Eckbo, B. E., editor, *Handbook of Corporate Finance: Empirical Corporate Finance*, volume 1, pages 3–36. North-Holland.

- Lee, D. S. (2008). Randomized Experiments from Non-Random Selection in U.S. House Elections. *Journal of Econometrics*, 142(2):675–697.
- Marshall, J. (2024). Can Close Election Regression Discontinuity Designs identify Effects of Winning Politician Characteristics? *American Journal of Political Science*, 68(2):494–510.
- Rainey, C. (2014). Arguing for a Negligible Effect. *American Journal of Political Science*, 58(4):1083–1091.
- Snyder Jr., J. M. (1992). Long-term Investing in Politicians; or, Give Early, Give Often. *Journal of Law and Economics*, 35(1):15–43.
- Wu, Z. and Cao, W. (2026). Competitive Dynamics in the Political Marketplace: How Donor Firms’ Political Wins Prompt Product-Market Peers to Appoint Politicians to Their Boards. *Management Science*, 72(4). Published online August 2025.
- Ziobrowski, A. J., Cheng, P., Boyd, J. W., and Ziobrowski, B. J. (2004). Abnormal Returns from the Common Stock Investments of the U.S. Senate. *Journal of Financial and Quantitative Analysis*, 39(4):661–676.