

Does Party Affiliation Trump Appointing President? An Empirical Analysis of Trump’s Appellate Judges

Jonathan P. Kastellec
Department of Politics
Princeton University
jkastell@princeton.edu

Eitan Sapiro-Gheiler
Department of Politics
Princeton University
eitans@princeton.edu

July 20, 2026

Abstract

We examine whether President Trump’s appointees to the U.S. Courts of Appeals act differently from other Republican appointees. Using a new dataset covering the universe of published U.S. Courts of Appeals opinions, we show that three-judge panels with Trump-appointed judges do not produce more conservative outcomes than panels with other Republican presidents’ appointees. Trump appointees are, however, more likely than non-Trump-Republican appointees to cast conservative votes from the bench—a gap that widens with the number of Democratic-appointed judges on the panel. This difference between panel outcomes and individual votes is driven by Trump appointees’ increased tendency to dissent from liberal majorities; they also write concurring opinions more often, especially when sitting with Democratic appointees. This evidence from the federal judiciary speaks to the broader debate over whether Trump represents a continuation of, versus a departure from, longstanding trends in Republican Party politics.

*AI disclosure: The Comprehensive Courts of Appeals Database (CCAD), the paper’s primary data source, was constructed using a large language model to code the ideological direction and disposition of appellate opinions following the coding conventions of the Songer Court of Appeals Database and the Supreme Court Database; this pipeline is documented in Sapiro-Gheiler and Kastellec (2026) and validated against those human-coded benchmarks. For the salient topic analyses in Online Appendix A.4.1 (gun rights, religious liberty, and immigration), we used a large language model to identify relevant cases and to code the claims at issue and their outcomes; details are in that section. These codings were validated against existing CCAD measures and keyword-based searches of the opinion text, and spot-checked by the authors. Finally, we used a large language model (Anthropic’s Claude Opus 4.7 via Claude Code) for copyediting and for identifying errors in prose and LaTeX formatting. The authors conducted all statistical analysis, wrote the substantive text of the manuscript, and are solely responsible for the validity of all results, interpretations, and any elements produced with AI assistance.

1 Introduction

President Donald Trump has reshaped the federal judiciary during his two terms in office. In addition to appointing three Supreme Court justices, as of the end of 2025 Trump had appointed about one-third of active judges on the Courts of Appeals and the district courts. While Trump’s continued ability to appoint judges to the federal bench will depend on Senate control in the 2027-28 Congress, his judicial legacy is already assured to be long-lasting.

As with many of Trump’s actions, a key empirical question is to what extent his judicial appointees represent a break with previous Republican presidents versus a continuation of partisan tendencies. Compared to his Republican predecessors, Trump’s nominees have been much more likely to be affiliated with the Federalist Society and the National Rifle Association; they are also more likely to have religious affiliations (Choi, Gulati and Posner 2025). Evidence on whether Trump-appointed judges vote distinctively from the bench is concentrated in high-salience issue areas corresponding to these biographical differences. On the Courts of Appeals, Brown, Epstein and Gulati (2025) find that Trump-appointed judges are more likely than non-Trump-Republican appointees to support gun rights in Second Amendment cases, while Rothschild (2022) and Choi, Gulati and Posner (2025) find greater support for Christian plaintiffs in free exercise cases. In a study restricted primarily to younger judges, Choi and Gulati (2025) find that Trump appointees are more likely to be highly cited and somewhat more likely to write opinions, but no more likely to write separate (i.e., dissenting or concurring) opinions. On the district courts, Manning, Carp and Holmes (2020) report that Trump judges vote more conservatively on average than any presidential cohort since the 1960s, though they present only raw means with no adjustment for case characteristics, while Boldt et al. (2026) find no partisan differences—among Trump appointees or otherwise—in January 6th prosecutions in the District of Columbia.

While useful, these studies are limited in scope and statistical power. Gun rights and religious freedom cases together account for less than 5% of appellate cases from 2017–2025. Even setting that sample’s representativeness aside, its size prevents prior work from leverag-

ing the Courts of Appeals’ quasi-random assignment of judges to panels. This process—the key to causally identified estimates of judicial behavior—occurs within circuits and years, so limited samples feature only a handful of cases where Trump-appointed judges sit with other Republican appointees. Below the Courts of Appeals, district judges sit alone and identification necessarily comes from between-case variation. Above them, Supreme Court justices decide cases as a multi-member court, but select which cases to review and always sit *en banc*, making full causal identification impossible. (The Supreme Court offers a further argument against Trump appointees being clearly more conservative: the most reliable conservative votes, Clarence Thomas and Samuel Alito, were *not* appointed by Trump.)

We examine whether Trump’s judges are systematically different from other Republican appointees across all appellate cases rather than within a particular issue area, using a new dataset—the [Comprehensive Courts of Appeals Database](#)—that covers the universe of published Courts of Appeals opinions from 1892–2025. Across over 30,000 cases decided from 2017–2025, we show that panels with and without Trump appointees produce the same share of conservative *case outcomes*, with that share mostly depending on the total number of Republican appointees—that is, the overall partisan composition of the panel, regardless of within-appointing-party differences by appointing president. Trump appointees are, however, more likely than non-Trump-Republican appointees to cast conservative *votes from the bench*—a pattern driven by their greater willingness to dissent from liberal majorities. This gap in conservative voting grows with the number of Democratic-appointed co-panelists, but even at its largest it remains smaller than the overall gap between Republican and Democratic appointees. Finally, we find that Trump appointees are also more likely than non-Trump-Republican appointees to write concurring (or partially concurring) opinions, especially when they sit with Democratic appointees. These results suggest that Trump’s appointees represent only a limited departure from longstanding trends in Republican Party politics (Barber and Pope 2019).

2 Data and Motivating Pattern

Our primary data source is the [Comprehensive Courts of Appeals Database](#) (CCAD), which combines metadata for the universe of published appellate cases with a large-language-model (LLM) powered extension of the research-standard Songer U.S. Courts of Appeals Database (Songer 2008). We consider only cases decided by standard three-judge panels—excluding en banc decisions and panels with fewer judges—in the twelve regional U.S. Courts of Appeals—including the D.C. Circuit and excluding the Federal Circuit. We also exclude fewer than 50 cases with more than one authored dissent.

For all main-text analyses involving the ideological direction (i.e., liberal, conservative, mixed, or not ascertained) of a case outcome or judge vote, we use the CCAD’s “Songer-track super-mechanical direction,” the best-performing measure as validated on the hand-coded Songer database (Sapiro-Gheiler and Kastellec 2026). Results are robust to using other ideological direction measures provided by the CCAD; see Appendix A.2.1 for details. As in the original Songer Database, the CCAD’s ideological direction measure (which uses an LLM to replicate the hand-coded Songer scheme) is determined by the case type and litigants. For example, a vote in support of a criminal defendant will always be coded as “liberal.”

Our analyses use cases from 2017–2025 to isolate differences beginning in Trump’s first term. However, to provide context, we document historical trends in outcomes by appellate panel composition, commonly called “panel effects,” over time (Kastellec 2011). For 1925–2025 (a sample containing around 380,000 cases meeting our criteria), we compute the average share of conservative case outcomes in each five-year period for each panel composition: all Republican appointees (RRR), two Republican appointees (RRD), two Democratic appointees (DDR), or all Democratic appointees (DDD). Results are shown in Figure 1.

Consistent with broad trends in partisan polarization, there was not much difference across panel compositions until the mid-20th century. However, beginning in the 1970s we see a striking divergence. Majority-Republican panels increasingly voted conservatively, while the share among Democratic-majority panels decreased before rebounding. In the 2020s,

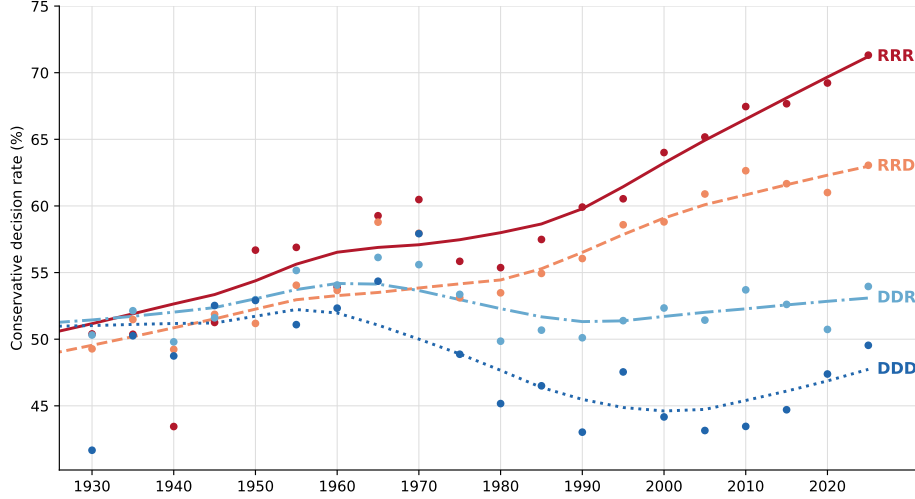


Figure 1: Conservative decision rate by panel composition, 1925–2025. Points are means in five-year bins; x -axis labels are bin endpoints. Lines are LOESS smooths (span 0.5).

the difference between RRR panels and DDD panels was over 20 percentage points, almost a 50% increase over the share of conservative outcomes on DDD panels. For our purposes, these results suggest that polarization between Democratic and Republican appointees well precedes the influx of Trump-appointed judges; simply regressing the share of conservative case outcomes from 1970–2025 (or even 2000–2025) on an indicator for Trump appointees would mechanically give a positive result, capturing the overall increasing trend in judicial partisanship during Trump’s terms in office. The true outcome of interest is whether Trump appointees are more conservative than their colleagues who serve at the same time or even on the same panels, which is what we test in the next section.

3 Case Outcomes, Judge Votes, and Separate Opinions

We now turn our focus to 2017–2025, when Trump-appointed judges sit on the bench. Our sample contains 32,771 standard three-judge panels with all judges and opinion authors identified and at most one dissenting opinion.¹ Table 1 reports the total number of authored opinions from 2017–2025, including panels we drop because of missing data.

¹The number of observations differs slightly for some specifications because some of these panels are in degenerate circuit-by-year cells or lack observations of control variables.

Type	Trump app.	Other R app.	D app.	Per curiam	Not id.	Total
Majority	3,388	14,882	12,592	445	74	31,381
Dissent	524	654	923	0	140	2,241
Concurrence	602	532	869	0	142	2,145
Mixed	200	252	376	0	20	848
Rehearing	163	723	781	55	1	1,723
<i>Total</i>	4,877	17,043	15,541	500	377	38,338

Table 1: All Courts of Appeals opinions by appointing president group, 2017–2025.

Our analysis covers the three important observable features of each panel’s decision: the ideological direction² of the panel’s *case outcome*, the ideological direction of each judge’s *vote* (a majority vote matches the case outcome’s direction, a dissenting vote is opposite), and the presence of *separate opinions* (dissenting, concurring, or both).

3.1 Case Outcomes

We begin our analysis of case outcomes with a basic descriptive question: does the share of conservative outcomes differ between panels with Trump-appointed judges and panels with non-Trump-Republican appointees? For any panel with at least one Republican-appointed judge, we condition on the number of Democratic-appointed co-panelists and compare the share of conservative case outcomes when *all* Republican appointees are Trump-appointed versus when *none* are (excluding panels with both). The results are presented in Table 2. Consistent with Figure 1, there are substantial decreases in the share of conservative case outcomes as the number of Democratic appointees increases (i.e., moving down the rows). By contrast, the differences between panels with and without Trump appointees, conditional on partisan composition, are about a quarter as large. Panels with Trump judges are actually slightly *less* likely to reach a conservative case outcome.

These aggregate statistics are suggestive, but not causally identified; they compare judges

²As noted above, the mechanical nature of coding ideological direction based on the Songer Database makes it possible to assign direction to a wide variety of cases, but it may also result in a lack of nuance for specific issue areas. In Section A.4.1, we show that the CCAD’s rich data allows a more tailored approach for topics of interest, where we find some Trump-appointee-specific differences.

Composition	All-Trump		No-Trump		Gap (pp)
	Cons. %	n	Cons. %	n	
RRR	68.6	338	71.6	2,410	-3.0
RRD	59.8	1,561	62.1	5,109	-2.3
DDR	50.3	2,792	52.8	6,322	-2.6

Table 2: Share of conservative case outcomes by panel composition and appointing president, 2017–2025. Rows fix partisan composition; columns compare panels whose Republican-appointed seats are all Trump-appointed versus none Trump-appointed (excluding panels with both).

across both circuits and time. To leverage quasi-random assignment of judges to panels and panels to cases, we regress the ideological direction of each case’s dispositional outcome on measures of panel partisanship using circuit-by-year fixed effects. For panel i in circuit c and year t , we estimate linear probability models of the form

$$y_i = \sum_{k=1}^3 \alpha_k R_i^{(k)} + \sum_{k=1}^3 \beta_k (R_i^{(k)} \cdot \# \text{Trump}R_i) + \gamma_{ct} + \mathbf{x}_i' \delta + \varepsilon_i,$$

where y_i is an indicator for a conservative case outcome; the $R_i^{(k)}$ are indicators for k total Republican appointees on the panel (zero Republican appointees is the omitted category); $\# \text{Trump}R_i$ counts the number of Trump appointees on the panel; γ_{ct} are circuit-by-year fixed effects; and $\mathbf{x}_i$ collects controls. The main effects α_k are the baseline level of conservative outcomes on panels with k total Republican appointees relative to omitted all-Democratic panels; the interactions β_k are the marginal effect of an additional Trump appointee (instead of a non-Trump-Republican appointee) on panels with k total Republican appointees.

The results from this specification are shown in Table 3. Models (1) and (2) include all panels, with Model (2) adding controls for the case’s broad issue area and panel demographics. Models (3) and (4) restrict the sample to panels of circuit judges only (i.e., without any district judges sitting by designation), adding the same controls in Model (4). As expected, increasing the number of Republican appointees on a panel increases the share of conservative outcomes—by about seven percentage points for each additional one. Trump appointees slightly decrease the chance of a conservative outcome on Republican-appointee-

	All panels		Circuit judges only	
	(1)	(2)	(3)	(4)
#Trump-R seat $\times$ RRR	-0.0024 (0.0076)	-0.0083 (0.0084)	-0.0022 (0.0079)	-0.0079 (0.0089)
#Trump-R seat $\times$ RRD	0.0026 (0.0073)	-0.0087 (0.0075)	0.0001 (0.0079)	-0.0104 (0.0081)
#Trump-R seat $\times$ DDR	0.0108 (0.0148)	0.0083 (0.0149)	0.0150 (0.0150)	0.0132 (0.0153)
RRR panel	0.2226*** (0.0166)	0.2302*** (0.0186)	0.2168*** (0.0184)	0.2235*** (0.0196)
RRD panel	0.1571*** (0.0145)	0.1664*** (0.0149)	0.1523*** (0.0159)	0.1598*** (0.0157)
DDR panel	0.0658*** (0.0129)	0.0690*** (0.0133)	0.0528*** (0.0131)	0.0562*** (0.0133)
DDD rate (modeled)	0.4606	0.4590	0.4724	0.4715
Circuit $\times$ year FE	✓	✓	✓	✓
Broad issue area FE		✓		✓
Panel controls		✓		✓
Observations	32,283	32,283	28,696	28,696
R ²	0.0513	0.1003	0.0516	0.1034

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

Standard errors in parentheses, clustered by circuit $\times$ year. Broad issue area fixed effects are from the CCAD. Panel controls are indicators for ≥ 1 woman and ≥ 1 nonwhite judge on the panel, as well as the panel's mean age.

The omitted category for all models is DDD panels.

Table 3: Share of conservative case outcomes, 2017–2025. Interactions track the Trump-appointee difference relative to non-Trump-Republican appointees, holding fixed the total number of Republican appointees on a panel; main effects are changes on panels without Trump appointees compared to the DDD panel baseline.

majority panels and slightly increase it on Democratic-appointee-majority panels, but no point estimates are statistically significant and even the largest are only around 1.5 percentage points—one-fifth of the per-seat between-appointing-party effect. Replacing a Trump appointee with a non-Trump-Republican appointee does not change the likelihood of a conservative case outcome, *ceteris paribus*.

3.2 Judge Votes

While the null results in the case outcome models of Table 3 are informative, they may still mask judge-level differences. If Trump appointees issue conservative dissents more often than

non-Trump-Republican appointees, the conservative outcome share would not be affected by these non-precedential opinions, but we might still consider Trump judges to be more extreme. We thus turn to analyzing judge votes.

Our first approach is similar to Table 3, switching the level of estimation to judge votes rather than cases. For judge j on panel i in circuit c and year t , we estimate linear probability models of the form

$$y_{ij} = \sum_{k=1}^2 \alpha_k R_i^{(k)} + \sum_{k=1}^3 \beta_k (R_i^{(k)} \cdot \text{Trump}R_j) + \sum_{k=1}^3 \rho_k (R_i^{(k)} \cdot \text{nonTrump}R_j) + \gamma_{ct} + \mathbf{x}_i' \delta + \varepsilon_{ij},$$

where y_{ij} is an indicator for a conservative vote on panel i by judge j ; the $R_i^{(k)}$ are indicators for k total Republican appointees on the panel (zero Republican appointees is the omitted category); $\text{Trump}R_j$ and $\text{nonTrump}R_j$ are indicators for whether judge j is, respectively, a Trump appointee or a non-Trump-Republican appointee (Democratic appointees are the omitted category); γ_{ct} are circuit-by-year fixed effects; and $\mathbf{x}_i$ collects controls. The baseline group is now Democratic appointees rather than non-Trump-Republican appointees. The main effects α_k therefore capture the conservative voting rate of Democratic appointees on panels with k total Republican appointees (relative to omitted all-Democratic panels); the first set of interactions β_k capture the difference between a Trump appointee and a Democratic appointee; and the second set ρ_k capture the difference between a non-Trump-Republican appointee and a Democratic appointee.³ This specification thus identifies differences in conservative voting rates by comparing judges between panels.

The results are shown in models (1) and (2) of Table 4; paralleling models (2) and (4) of Table 3, the former includes all panels while the latter restricts to circuit-judge-only panels (both including controls). For simplicity, we show only the $\text{Trump}R - \text{nonTrump}R$ coefficient difference and the modeled conservative voting rate for non-Trump-Republican appointees. On Republican-appointee-majority (RRR and RRD) panels, there is no statistically significant difference between Trump appointees and non-Trump-Republican appointees. However,

³Because no Democratic appointee ever sits on an all-Republican panel, the omitted group on those panels is non-Trump-Republican appointees: α_3 drops out, β_3 compares Trump appointees to non-Trump-Republican appointees, and ρ_3 is the non-Trump-Republican-appointee baseline.

	Between panels		Within panels	
	All panels (1)	Circuit judges only (2)	All panels (3)	Circuit judges only (4)
<i>Trump-R – non-Trump-R</i>				
RRR	-0.0049 (0.0082)	-0.0055 (0.0083)	-0.0005 (0.0025)	-0.0010 (0.0025)
RRD	0.0055 (0.0099)	0.0043 (0.0101)	0.0172*** (0.0045)	0.0173*** (0.0047)
DDR	0.0438** (0.0149)	0.0417** (0.0153)	0.0322** (0.0098)	0.0327** (0.0100)
<i>non-Trump-R rate (modeled)</i>				
RRR	0.6848	0.6874	0.7013	0.7026
RRD	0.6272	0.6303	0.6179	0.6195
DDR	0.5554	0.5601	0.5492	0.5514
Case FE			✓	✓
Circuit × year + issue FE	✓	✓		
Panel controls	✓	✓		
Observations	96,849	93,262	96,849	93,262
R ²	0.0957	0.0967	0.9186	0.9197

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

Standard errors in parentheses, two-way clustered by case and judge. Broad issue area fixed effects are from the CCAD; they are absorbed by case fixed effects in models (3) and (4). Panel controls are indicators for ≥ 1 woman on the panel, and ≥ 1 nonwhite judge on the panel, as well as the panel's mean age; they are absorbed by case fixed effects in models (3) and (4).

F-tests for differences between the RRD and DDR coefficients have *p*-values of 0.011, 0.015, 0.123, and 0.118 for models (1)–(4) respectively.

Table 4: Conservative votes by judge, 2017–2025. Models (1) and (2) compare judges between panels controlling for case and panel features; models (3) and (4) compare judges on the same panel.

on panels with only one Republican appointee (DDR panels), Trump appointees are about four percentage points more likely to vote conservatively—just over half the size of the modeled difference between RRD and DDR panels. Given the lack of Trump-specific differences in conservative case outcomes, this difference in votes suggests that Trump appointees dissent more often from liberal majorities; we return to this point in the next section.

While the between-case specification in models (1) and (2) is causally identified, it does not fully account for idiosyncratic variation in case facts and panels. To sharpen our comparison between Republican appointing presidents, we define a specification with case fixed effects, exploiting quasi-random assignment and controlling for panel-level differences by comparing judges within the same panel. We estimate a linear probability model of the form

$$y_{ij} = \sum_{k=1}^3 \beta_k (R_i^{(k)} \cdot \text{TrumpR}_j) + \sum_{k=1}^3 \rho_k (R_i^{(k)} \cdot \text{nonTrumpR}_j) + \kappa_i + \varepsilon_{ij},$$

where, compared to the between-case specification, the case fixed effects κ_i absorb that specification’s main effects α_k , the circuit-by-year fixed effects γ_{ct} , and the controls $\mathbf{x}_i$. In this specification, the β_k and ρ_k are identified using differences in votes on the same panel, rather than voting rates across all panels with k Republican appointees. Variation thus depends on the presence of a dissent, which is relatively rare on the Courts of Appeals—there are only about 500 identifying panels in our data where Trump appointees and non-Trump-Republican appointees vote differently. However, the specification uses the contrast with *Democratic* appointees—about 2,000 more identifying panels—to obtain tighter estimates and impute the $\text{TrumpR} - \text{nonTrumpR}$ comparison on DDR panels.

The results for this specification are shown in models (3) and (4) of Table 4, again considering all panels and circuit-judge-only panels. The between-case results are replicated on RRR and DDR panels, albeit with a one-percentage-point smaller (but still large and statistically significant) effect in the latter case. However, on RRD panels, we now find a larger and statistically significant increase in conservative votes by Trump appointees—around 1.7 percentage points. This difference suggests that idiosyncratic variation not fully captured by case and panel controls may have dampened the between-panel estimate.

3.3 Separate Opinions: Dissents and Concurrences

Because Trump’s judges do not produce more conservative *case outcomes*, but do cast more conservative *votes*, they must be writing conservative dissents more often than other Republican appointees. In this section, we verify that fact and also test whether Trump appointees are more likely to write concurring separate opinions as well. Separate opinions, by definition, do not carry the same weight as majority opinions, but they often lay groundwork for amending or overturning prior rulings; concurring opinions also offer alternative rationales or rules that other courts may choose to follow.

	Between panels			Within panels		
	Dissent	Concurrence	Mixed	Dissent	Concurrence	Mixed
<i>Trump-R – non-Trump-R</i>						
RRR	-0.0004 (0.0026)	0.0060 (0.0042)	0.0027 (0.0015)	-0.0005 (0.0025)	0.0068 (0.0039)	0.0028 (0.0015)
RRD	0.0085* (0.0041)	0.0131** (0.0043)	0.0020 (0.0016)	0.0044 (0.0040)	0.0141*** (0.0037)	0.0041* (0.0018)
DDR	0.0445*** (0.0108)	0.0350*** (0.0076)	0.0146* (0.0062)	0.0470*** (0.0092)	0.0375*** (0.0069)	0.0133* (0.0052)
<i>non-Trump-R rate (modeled)</i>						
RRR	0.0161	0.0142	0.0046	0.0119	0.0112	0.0036
RRD	0.0172	0.0157	0.0055	0.0175	0.0147	0.0045
DDR	0.0337	0.0184	0.0151	0.0355	0.0199	0.0161
Circuit × year + issue FE	✓	✓	✓			
Panel controls	✓	✓	✓			
Case FE				✓	✓	✓
Observations	96,849	96,849	96,849	96,849	96,849	96,849
R^2	0.0189	0.0124	0.0089	0.3259	0.3224	0.3305

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

Standard errors in parentheses, two-way clustered by case and judge. The between-panel model uses circuit-by-year fixed effects, broad issue area fixed effects from the CCAD, and panel controls (indicators for ≥ 1 woman on the panel, and ≥ 1 nonwhite judge on the panel, as well as the panel's mean age). The within-panel model uses case fixed effects, which absorb all of the between-panel model's fixed effects and controls.

Table 5: Separate opinion writing, 2017–2025. The between-panel model follows model (1) from Table 4; the within-panel model follows model (3) from that table.

We reuse the between-panel design of Model (1) from Table 4 (all panels; circuit-by-year fixed effects, broad issue area fixed effects, and panel controls) and the within-panel design of Model (3) from that table (all panels; case fixed effects absorbing all prior fixed effects and controls), now with the probability of writing a separate opinion as the dependent variable. Table 5 presents the results. As the vote-level results suggested, Trump appointees have a larger separate-writing gap to non-Trump-Republican appointees as the number of Democratic-appointed co-panelists increases. All specifications give similar results. Trump appointees do not write more separate opinions on RRR panels; on RRD panels, Trump appointees write about 0.5 percentage points more dissents, 1.5 percentage points more concurrences, and 0.25 percentage points more mixed (concurring-in-part-and-dissenting-in-part) opinions; on DDR panels, the gap is about 4.5, 3.5, and 1.5 percentage

points respectively. We do not model separate opinions’ ideological direction (e.g., liberal dissents from conservative majorities vs. conservative dissents from liberal majorities) since the decision to write separately interacts with a case outcome’s ideological direction. Descriptively, though, that pattern is informative: Trump appointees write more conservative dissents, fewer liberal dissents, and more of both liberal and conservative concurrences.

The headline result that replacing a non-Trump-Republican appointee with a Trump appointee produces no change in the share of conservative case outcomes thus masks a consistent and meaningful difference in non-precedential votes from the bench. Trump appointees are more likely than non-Trump-Republican appointees to write separately, especially in dissent from liberal majorities, resulting in a higher conservative voting rate. This tendency is heightened on panels with more Democratic appointees, where Republican appointees as a group have less control over the panel’s decision. Whether this higher rate of conservative separate opinions will eventually spill over into case outcomes, or affect the stability of legal precedent overall, is an important question for future work.

3.4 Robustness Checks and Extensions

Before concluding, we briefly summarize the robustness checks and extensions presented in the Online Appendix. We first confirm that the motivating pattern of increasing polarization over time holds in a causally identified model with circuit-by-year fixed effects (Section A.1). Next, we verify that our main results are robust to alternative ideological direction codings from the CCAD (Section A.2.1), separating effects by broad issue area (Section A.2.2), or dropping judges who write many separate opinions (Section A.2.3). We discuss statistical power (minimum detectable effects and test calibration) in Section A.3; despite some overly narrow confidence intervals for DDR panels in particular, our conclusions are not affected.

Finally, we provide two substantive extensions. Section A.4.1 applies our approach to three particularly salient topics for Trump appointees: gun rights, religious liberty, and immigration. Trump appointees are generally more supportive of gun-rights and Christian

claimants, and less supportive of immigration claimants/noncitizens, though the differences are not always statistically significant and depend on how the outcome is coded. Section A.4.2 presents parallel president-specific models for the appointees of Presidents George W. Bush and Joseph R. Biden; there is no distinctive pattern in case outcomes or judge votes for either president’s appointees.

4 Conclusion

Using new data covering the universe of published U.S. Courts of Appeals cases, we showed that panels including President Trump’s appointees do not reach conservative case outcomes more often than those including other Republican presidents’ appointees. However, Trump appointees vote more conservatively from the bench than their Republican-appointed colleagues, dissenting more often from liberal majorities and less often from conservative ones. When joining the majority, Trump appointees are more likely to stake out an individual position in a concurring opinion. These results deepen our understanding of broad patterns in appellate cases, adding to a literature that has focused on salient topics rather than aggregate outcomes. They also shed light on the extent to which Trump’s judicial nominations and their policy effects represent a departure from Republican Party politics.

Our results do not rule out the possibility that the presence of Trump appointees pushes appellate panels to rule more conservatively on particular cases, and we do not analyze the additional content conveyed by Trump appointees’ increased volume of separate opinions. Additionally, the subtle differences we find may well propagate as Trump continues to appoint more judges and as older non-Trump-Republican appointees exit the bench. Future work in this area could continue to build upon our approach of large-scale statistical analysis with tools from natural language processing and qualitative studies of salient cases to strengthen our understanding of judicial polarization.

References

- Barber, Michael and Jeremy C. Pope. 2019. “Does Party Trump Ideology? Disentangling Party and Ideology in America.” *American Political Science Review* 113(1):38–54.
- Boldt, Ethan D., Hayley Munir, Jason S. Byers and Michael C. Gizzi. 2026. ““Stop the Steal” Sentenced: An Analysis of Judge Appointment Effects on the Sentencing of January 6 Defendants.” *Political Research Quarterly* .
- Brown, Rebecca L., Lee Epstein and Mitu Gulati. 2025. “Guns, Judges, and Trump.” *Duke Law Journal Online* 74:81–109.
- Choi, Stephen J and Mitu Gulati. 2025. “How Different Are the Trump Judges?” *Vanderbilt Law Review En Banc* 78(1):1–62.
- Choi, Stephen J., Mitu Gulati and Eric A. Posner. 2025. “Trump’s Lower Court Judges and Religion: An Initial Appraisal.” *The Journal of Legal Studies* 54(1):1–41.
- Kastellec, Jonathan P. 2011. “Panel Composition and Voting on the U.S. Courts of Appeals Over Time.” *Political Research Quarterly* 64(2):377–391.
- Manning, Kenneth L., Robert A. Carp and Lisa M. Holmes. 2020. “The Decision-Making Ideology of Federal Judges Appointed by President Trump.” SSRN working paper, available at https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3716378.
- Rothschild, Zalman. 2022. “Free Exercise Partisanship.” *Cornell Law Review* 107:1067–1136.
- Sapiro-Gheiler, Eitan and Jonathan P. Kastellec. 2026. “Appealing to Large-Language-Model-as-Judge: A Comprehensive Machine-Coded Database for the U.S. Courts of Appeals.” Working paper, available at coadata.org.
- Songer, Donald R. 2008. “The United States Courts of Appeals Database.” <http://www.songerproject.org/us-courts-of-appeals-databases.html>. Original U.S. Courts of Appeals database, covering 1925–1996. Data last updated 21 October 2008.

Online Appendix

A.1	Polarization Over Time: Formal Results	15
A.2	Robustness Checks	17
A.2.1	Alternative Ideological Direction Codings	17
A.2.2	Issue Area Heterogeneity	18
A.2.3	Dropping Prolific Separate Writers	21
A.3	Statistical Power	22
A.4	Extensions	26
A.4.1	Salient Topics	26
A.4.2	Other Presidents' Appointees	33

A.1 Polarization Over Time: Formal Results

Figure 1 presented descriptive differences in conservative outcome rates across panel types; Table A.1 confirms these results formally. Model (1) uses circuit effects and a linear decade trend, while Model (2) uses circuit-by-year fixed effects (the level of quasi-random assignment of judges to panels and panels to cases). We restrict this analysis to cases decided in or after 1970, the point at which the divergence between panel types emerges.

	Trend Circuit FE (1)	Causal Circuit $\times$ year FE (2)
RRR	0.0476** (0.0152)	0.1190*** (0.0067)
RRD	0.0229 (0.0141)	0.0875*** (0.0060)
DDR	0.0065 (0.0119)	0.0396*** (0.0054)
RRR $\times$ Decade	0.0305*** (0.0045)	—
RRD $\times$ Decade	0.0273*** (0.0042)	—
DDR $\times$ Decade	0.0138*** (0.0036)	—
Decade (since 1970)	-0.0131** (0.0041)	—
Circuit FE	✓	
Circuit $\times$ year FE		✓
Observations	279,730	279,730
R ²	0.0215	0.0290

* $p < 0.05$ ** $p < 0.01$ *** $p < 0.001$
Standard errors in parentheses, clustered by circuit $\times$ year.
The omitted category is DDD panels.

Table A.1: Panel effects over time. Model (1) uses circuit fixed effects and a decade trend. Model (2) uses circuit-by-year fixed effects for causal identification, but in doing so absorbs the time trend.

In Model (1), the conservative-decision gap relative to all-Democratic panels widens by 1.38, 2.73, and 3.05 percentage points per decade for DDR, RRD, and RRR panels respectively. The implied DDD-to-RRR gap therefore grows from about 5 percentage points in 1970 to about 20 percentage points by 2025. The decade main effect is roughly 1 percentage point, so DDD panels drifted slightly leftward, while the absolute increase in conservative outcomes on RRR panels is about 2 percentage points per decade. In the causally identified Model (2), the gaps to DDR, RRD, and RRR panels (pooled across the whole time period) are, as expected, about half of the 2025 gap implied by the linear trend in Model (1).

A.2 Robustness Checks

This appendix section collects robustness checks for the three main analyses in the paper—case outcomes, judge votes, and separate opinion writing.

A.2.1 Alternative Ideological Direction Codings

All results in the main text use the CCAD’s Songer-track super-mechanical ideological direction measure. The CCAD also provides five other codings, formed by crossing two ground-truth tracks—Songer and the Supreme Court Database (SCDB)—with three derivations: super-mechanical, mechanical measure, and one-shot (Sapiro-Gheiler and Kastellec 2026). All six use the Songer-standard 0/1/2/3 direction scale (0 = not ascertained, 1 = conservative, 2 = mixed, 3 = liberal) but SCDB-track codings never produce the mixed code (which is not in the SCDB codebook). To confirm that our results do not depend on our choice of coding, we re-estimate our main results from Tables 3 and 4 using each of them.

For ease of comparison, the result shown in Tables A.2 and A.3 place the **TrumpR** vs. **nonTrumpR** difference on each panel type (RRR, RRD, and DDR) in its own column, with the rows showing its value under each ideological direction coding. Both tables broadly match the main text; case outcome differences are around one percentage point in magnitude and not statistically significant, while vote-level differences become larger as the number of Democratic appointees on the panel increases, with statistically significant effects for all between-panel DDR and within-panel RRD and DDR effects except when using Songer one-shot direction (which produces no statistically significant effects at all).

Direction coding	All panels				Circuit judges only			
	<i>N</i>	RRR	RRD	DDR	<i>N</i>	RRR	RRD	DDR
Songer super-mech. [†]	32,283	-0.0083 (0.0084)	-0.0087 (0.0075)	0.0083 (0.0149)	28,696	-0.0079 (0.0089)	-0.0104 (0.0081)	0.0132 (0.0153)
Songer mech.	32,237	-0.0050 (0.0093)	-0.0081 (0.0089)	0.0128 (0.0147)	28,657	-0.0050 (0.0098)	-0.0106 (0.0094)	0.0170 (0.0155)
Songer one-shot	32,198	-0.0053 (0.0083)	-0.0062 (0.0092)	0.0165 (0.0147)	28,648	-0.0066 (0.0087)	-0.0088 (0.0095)	0.0173 (0.0151)
SCDB super-mech.	32,763	-0.0083 (0.0099)	-0.0087 (0.0096)	0.0098 (0.0150)	29,134	-0.0100 (0.0104)	-0.0128 (0.0101)	0.0170 (0.0164)
SCDB mech.	32,763	-0.0081 (0.0099)	0.0033 (0.0101)	0.0112 (0.0143)	29,134	-0.0103 (0.0103)	-0.0038 (0.0103)	0.0137 (0.0160)
SCDB one-shot	32,680	-0.0098 (0.0089)	-0.0004 (0.0087)	-0.0044 (0.0139)	29,060	-0.0100 (0.0093)	-0.0061 (0.0091)	-0.0023 (0.0156)

[†] Headline coding used in the main text.

Table A.2: Robustness of case outcome result to alternative ideological direction codings. The first column block re-estimates Model (2) from Table 3, showing the # Trump-R-seat coefficient for each panel type. The second column block repeats that exercise for Model (4).

Direction coding	Between panels				Within panels			
	<i>N</i>	RRR	RRD	DDR	<i>N</i> (ID)	RRR	RRD	DDR
Songer super-mech. [†]	96,849	-0.0049 (0.0082)	0.0055 (0.0099)	0.0438** (0.0149)	96,849 (2,888)	-0.0005 (0.0025)	0.0172*** (0.0045)	0.0322** (0.0098)
Songer mech.	96,711	-0.0036 (0.0078)	0.0044 (0.0101)	0.0416** (0.0140)	96,711 (2,885)	-0.0013 (0.0023)	0.0161*** (0.0045)	0.0268** (0.0096)
Songer one-shot	96,594	-0.0061 (0.0075)	-0.0030 (0.0082)	0.0180 (0.0137)	96,594 (2,853)	-0.0025 (0.0021)	0.0052 (0.0038)	-0.0004 (0.0090)
SCDB super-mech.	98,289	-0.0088 (0.0082)	0.0021 (0.0098)	0.0464** (0.0148)	98,289 (2,915)	-0.0039 (0.0022)	0.0142** (0.0044)	0.0317** (0.0106)
SCDB mech.	98,289	-0.0075 (0.0078)	0.0100 (0.0103)	0.0352* (0.0137)	98,289 (2,915)	-0.0002 (0.0023)	0.0118** (0.0043)	0.0206* (0.0103)
SCDB one-shot	98,040	-0.0070 (0.0068)	0.0105 (0.0095)	0.0368** (0.0130)	98,040 (2,910)	0.0002 (0.0024)	0.0125** (0.0042)	0.0371*** (0.0101)

[†] Headline coding used in the main text.

Table A.3: Robustness of judge-vote result to alternative ideological direction codings. The “Between panels” block re-estimates Model (1) from Table 4, showing the **TrumpR** – **nonTrumpR** difference for each panel type. The “Within panels” block repeats that exercise for Model (3). “*N*” is the number of judge-vote observations for each model, with “(ID)” for the within-panel model showing the number of cases with identifying variation (differing votes across appointing president groups).

A.2.2 Issue Area Heterogeneity

As noted in the introduction and discussed in Section A.4.1, some prior work has found that Trump judges vote differently than other Republican appointees in specific topic areas

(Brown, Epstein and Gulati 2025, Rothschild 2022, Choi, Gulati and Posner 2025). To complement that section’s fine-grained analysis, we also re-estimate our main specifications from Tables 3 and 4 with added interactions for seven of the eight Songer broad issue areas; we drop Privacy, which accounts for less than 50 cases in our sample.

Case outcome results are shown in Table A.4. Unsurprisingly given the main text’s null result (Table 3), there are no statistically significant Trump effects on any panel type for any issue area. The broad pattern of negative point estimates for majority-R panels and positive ones for majority-D panels generally holds, but is more noisy due to sample size.

Vote-level results are shown in Table A.5. They once again match the main text (Table 4), with statistically significant results for between-panel DDR, within-panel DDR, and within-panel RRD estimates only. Among these, statistical significance appears to be driven mostly by sample size, appearing in the dominant Criminal, Civil Rights, and Economic Activity issue areas. The broad pattern of no difference on RRR panels, a small difference on RRD panels, and a large difference on DDR panels is mostly driven by these issue areas; it is also present for Due Process, and fully reverses for Miscellaneous (though the point estimates are not statistically significant for either of those issue areas).

	<i>N</i>	Marginal Trump-seat effect		
		RRR	RRD	DDR
Criminal	11,054	-0.0077 (0.0104)	-0.0174 (0.0133)	0.0145 (0.0247)
Civil rights	8,312	-0.0208 (0.0167)	0.0114 (0.0190)	-0.0319 (0.0277)
First Amendment	1,151	0.0009 (0.0540)	-0.0734 (0.0570)	0.0713 (0.0564)
Due process	431	-0.0386 (0.0955)	-0.0214 (0.1029)	-0.0170 (0.1198)
Labor relations	745	-0.0360 (0.0659)	-0.1029 (0.0610)	-0.0435 (0.0839)
Economic activity	9,883	0.0004 (0.0135)	0.0019 (0.0138)	0.0284 (0.0262)
Miscellaneous	681	0.0289 (0.0751)	0.0112 (0.0504)	0.1282 (0.1034)

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

Table A.4: Differences in conservative case outcomes separated by broad issue area, 2017–2025. We re-estimate Model (2) from Table 3 but interact the # Trump-R-seat coefficient with seven of the eight Songer broad issue areas (Privacy is omitted due to low sample size) from the CCAD. “*N*” is the number of cases for each broad issue area.

	Between panels				Within panels			
	<i>N</i>	RRR	RRD	DDR	<i>N</i> (ID)	RRR	RRD	DDR
Criminal	33,162	-0.0054 (0.0086)	-0.0019 (0.0124)	0.0606* (0.0254)	33,162 (806)	-0.0040 (0.0027)	0.0176** (0.0060)	0.0396** (0.0147)
Civil rights	24,936	-0.0114 (0.0138)	0.0267 (0.0178)	0.0282 (0.0269)	24,936 (923)	0.0023 (0.0050)	0.0224** (0.0077)	0.0525** (0.0194)
First Amendment	3,453	0.0036 (0.0315)	-0.0331 (0.0375)	0.0510 (0.0498)	3,453 (171)	0.0051 (0.0210)	0.0274 (0.0228)	-0.0367 (0.0406)
Due process	1,293	-0.0273 (0.0600)	0.0026 (0.0460)	0.0472 (0.0649)	1,293 (39)	0.0000 (0.0321)	0.0050 (0.0299)	0.0649 (0.0444)
Labor relations	2,235	-0.0043 (0.0460)	-0.0107 (0.0521)	0.0281 (0.0612)	2,235 (94)	0.0200 (0.0191)	0.0516 (0.0382)	0.0412 (0.0495)
Economic activity	29,649	-0.0042 (0.0113)	0.0064 (0.0136)	0.0510* (0.0204)	29,649 (744)	-0.0036 (0.0050)	0.0094 (0.0056)	0.0270* (0.0117)
Miscellaneous	2,043	0.0626 (0.0443)	0.0042 (0.0394)	-0.0210 (0.0699)	2,043 (107)	0.0412 (0.0354)	0.0109 (0.0254)	-0.1032 (0.0564)

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

Table A.5: Differences in judge-level conservative votes separated by broad issue area, 2017–2025. The “Between panels” block re-estimates Model (1) from Table 4 with interactions between the TrumpR and nonTrumpR coefficients and seven of the eight Songer broad issue areas (Privacy is omitted due to low sample size) from the CCAD. We show the TrumpR – nonTrumpR difference for each panel type. The “Within panels” block repeats that exercise for Model (3). “*N*” is the number of judge-vote observations for each model, with “(ID)” for the within-panel model showing the number of cases with identifying variation (differing votes across appointing president groups).

A.2.3 Dropping Prolific Separate Writers

Section 3.3 found that Trump appointees write separate opinions at higher rates than other Republican appointees. Because separate-opinion-writing is rare and concentrated among a few prolific authors, we test whether these gaps merely reflect the behavior of those outliers by re-estimating Table 5 after dropping the most prolific 5%, 10%, and 25% judges for each separate opinion type (among judges with more than 20 votes in our sample). Results in Table A.6 show that the effect on dissents is most robust; it attenuates to around half its magnitude after dropping the top 25% but remains statistically significant. The effect on concurrences generally survives only the 5 and 10% drops, and the effect on mixed opinions generally survives only the 5% drop. The pattern of larger effects on DDR panels than RRD panels is consistent at every drop level: the moderating influence of fellow Republican appointees affects the whole cohort.

	Between panels				Within panels			
	<i>N</i>	RRR	RRD	DDR	<i>N</i> (ID)	RRR	RRD	DDR
<i>Dissent</i>								
Full sample	96,849	-0.0004 (0.0026)	0.0085* (0.0041)	0.0445*** (0.0108)	96,849 (2,098)	-0.0005 (0.0025)	0.0044 (0.0040)	0.0470*** (0.0092)
Drop top 5% (18)	92,390	-0.0020 (0.0026)	0.0027 (0.0031)	0.0334*** (0.0097)	92,390 (1,683)	-0.0015 (0.0025)	-0.0017 (0.0037)	0.0343*** (0.0085)
Drop top 10% (35)	89,116	-0.0016 (0.0026)	0.0042 (0.0031)	0.0226* (0.0093)	89,116 (1,457)	-0.0015 (0.0025)	0.0019 (0.0035)	0.0231** (0.0081)
Drop top 25% (86)	77,793	0.0008 (0.0023)	0.0065* (0.0030)	0.0182* (0.0079)	77,793 (874)	0.0014 (0.0020)	0.0044 (0.0034)	0.0187** (0.0072)
<i>Concurrence</i>								
Full sample	96,849	0.0060 (0.0042)	0.0131** (0.0043)	0.0350*** (0.0076)	96,849 (1,891)	0.0068 (0.0039)	0.0141*** (0.0037)	0.0375*** (0.0069)
Drop top 5% (18)	92,771	0.0022 (0.0040)	0.0077* (0.0036)	0.0241*** (0.0063)	92,771 (1,477)	0.0031 (0.0039)	0.0104** (0.0033)	0.0249*** (0.0060)
Drop top 10% (35)	87,702	-0.0031 (0.0024)	0.0022 (0.0031)	0.0162** (0.0062)	87,702 (1,180)	-0.0038 (0.0022)	0.0038 (0.0028)	0.0172** (0.0059)
Drop top 25% (86)	77,171	-0.0022 (0.0023)	0.0001 (0.0023)	0.0080 (0.0049)	77,171 (730)	-0.0031 (0.0022)	0.0018 (0.0025)	0.0088 (0.0050)
<i>Mixed</i>								
Full sample	96,849	0.0027 (0.0015)	0.0020 (0.0016)	0.0146* (0.0062)	96,849 (794)	0.0028 (0.0015)	0.0041* (0.0018)	0.0133* (0.0052)
Drop top 5% (18)	94,620	0.0029* (0.0014)	0.0027 (0.0015)	0.0110* (0.0049)	94,620 (683)	0.0028 (0.0015)	0.0049** (0.0017)	0.0097* (0.0045)
Drop top 10% (35)	90,417	0.0024 (0.0014)	0.0012 (0.0013)	0.0054 (0.0046)	90,417 (555)	0.0019 (0.0015)	0.0032* (0.0016)	0.0040 (0.0043)
Drop top 25% (86)	76,675	0.0019* (0.0009)	0.0008 (0.0011)	0.0061 (0.0036)	76,675 (269)	0.0010 (0.0013)	0.0029* (0.0012)	0.0043 (0.0036)

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

Table A.6: Separate opinion writing dropping prolific judges. The “Between panels” block re-estimates Model (1) from Table 4, showing the TrumpR – nonTrumpR difference for each panel type. The “Within panels” block repeats that exercise for Model (3). “*N*” is the number of judge-vote observations for each model, with “(ID)” for the within-panel model showing the number of cases with identifying variation (different writing across appointing president groups).

A.3 Statistical Power

The meaningful difference in conservative votes by Trump appointees without any comparable effect on conservative outcomes raises questions about statistical power. To partially address these concerns, in Tables A.7 and A.8, we re-estimate the models in Tables 3 and 4 with all panel types pooled for additional statistical power. A single coefficient now captures the marginal effect of placing a Trump appointee on a panel instead of a non-Trump-Republican appointee, with the total number of Republican appointees held fixed.⁴

Results for conservative outcomes, following Table 3, are shown in Table A.7. In the main text’s headline Model (1), the minimum detectable effects (MDE) at 95% power were 3.04, 2.70, and 5.36 percentage points for RRR, RRD, and DDR panels respectively. The corresponding pooled specification produces an MDE of 2.03 percentage points—less than one-third of the difference between Republican and Democratic appointees overall—ruling out any total Trump-appointee difference (across all panels combined) larger than that amount.

In Table A.8, we show results for conservative votes, following Table 4. In the main text’s headline between-case Model (1), the MDEs at 95% power were 2.97, 3.58, and 5.37 percentage points for RRR, RRD, and DDR panels respectively. In the headline within-case Model (3), they were 0.89, 1.64, and 3.53 percentage points respectively. The corresponding pooled specifications produce MDEs of 2.66 and 1.15 percentage points respectively. The latter in particular essentially rules out any large difference in conservative voting by Trump appointees (while detecting a small one), showing the strength of the within-case comparison made possible by universe-level CCAD data.

⁴Formally, we replace the interaction terms in the case-outcome specifications with a single $\#TrumpR_i$ count variable and replace them in the judge-vote specifications with a $TrumpR_i$ indicator, but preserve the main effects to allow conservative outcome and vote rates to vary with the total number of Republican appointees on the panel.

	All panels		Circuit judges only	
	(1)	(2)	(3)	(4)
<i>Marginal Trump-appointed seat</i>				
All	0.0013 (0.0056)	-0.0065 (0.0065)	0.0008 (0.0060)	-0.0064 (0.0071)
<i>MDE (95% power, pp)</i>				
All	2.03	2.34	2.17	2.54
Circuit $\times$ year FE	✓	✓	✓	✓
Broad issue area FE		✓		✓
Panel controls		✓		✓
Observations	32,283	32,283	28,696	28,696

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

Table A.7: Differences in conservative case outcomes, re-estimating the models from Table 3 with the Trump-appointee effect pooled across all panel types to increase statistical power.

	Between panels		Within panels	
	All panels (1)	Circuit judges only (2)	All panels (3)	Circuit judges only (4)
<i>Trump-R – non-Trump-R</i>				
All	0.0064 (0.0074)	0.0052 (0.0075)	0.0115*** (0.0032)	0.0112*** (0.0032)
<i>non-Trump-R rate (modeled)</i>				
All	0.6092	0.6140	0.6075	0.6108
<i>MDE (95% power, pp)</i>				
All	2.66	2.71	1.15	1.17
Case FE			✓	✓
Circuit $\times$ year + issue FE	✓	✓		
Panel controls	✓	✓		
Observations	96,849	93,262	96,849	93,262
R ²	0.0956	0.0965	0.9186	0.9196

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

Table A.8: Judge-level differences in conservative votes, re-estimating the models from Table 4 with the Trump-appointee effect pooled across all panel types to increase statistical power.

The MDE computations above rely on well-calibrated confidence intervals. To confirm calibration, we perform a randomization inference exercise. Our standard errors are clustered at the judge and case level, so permuting labels within panels is the Fisher-exact test of our null hypothesis. However, that permutation would only change estimates for RRR and

RRD panels in the within-case model. We report results of the within-panel permutation for those coefficients only, and for the rest we permute at the circuit-by-year level. That is, we randomly re-assign the Trump and non-Trump labels within each circuit-by-year cell (meaning the same judge can have different labels in different cells). Each of the three permutations is performed 1,000 times, refitting our main models on each permutation. To compute p -values, we compare the t -statistic to its observed distribution, ensuring validity under our clustered standard errors (Chung and Romano 2013).

The results are shown in Tables A.9, A.10, and A.11–A.12, which provide the main-text and permutation p -values for the coefficients in Tables 3, 4, and 5 respectively. Each shows the main text p -value first, followed by the p -value from the circuit-by-year level permutation, followed by the p -value from the within-case permutation where feasible. Overall, p -values are generally comparable in size across the three methods. No formerly null effect is now statistically significant, and only a single coefficient that was statistically significant at the 5% level—the between-case estimate of increased dissents on RRD panels—is no longer statistically significant (the p -value rises from 0.038 to 0.055, marginal either way). We take these results as strong support for both our null results on case outcomes and our statistically and practically significant results on judge votes and separate opinions.

	Marginal Trump-R seat
#Trump-R seat $\times$ RRR	-0.0083 (0.325/0.267/—)
#Trump-R seat $\times$ RRD	-0.0087 (0.247/0.265/—)
#Trump-R seat $\times$ DDR	0.0083 (0.578/0.570/—)
Circuit $\times$ year FE	✓
Broad issue area FE	✓
Panel controls	✓
Observations	32,283
R ²	0.1003

Table A.9: Main text and permutation p -values for Trump appointees’ effects on case outcomes, following Model (2) of Table 3.

	Between panels	Within panels
<i>Trump-R – non-Trump-R</i>		
RRR	-0.0049 (0.550/0.528/—)	-0.0005 (0.824/0.822/0.843)
RRD	0.0055 (0.581/0.586/—)	0.0172 (<0.001/<0.001/0.002)
DDR	0.0438 (0.003/0.006/—)	0.0322 (0.001/0.006/—)
Case FE		✓
Circuit × year + issue FE	✓	
Panel controls	✓	
Observations	96,849	96,849
R ²	0.0957	0.9186

Table A.10: Main text and permutation p -values for Trump appointees’ effects on judge votes, following models (1) and (3) of Table 4.

	Dissent	Concurrence	Mixed
<i>Trump-R – non-Trump-R</i>			
RRR	-0.0004 (0.885/0.875/—)	0.0060 (0.152/0.139/—)	0.0027 (0.065/0.062/—)
RRD	0.0085 (0.038/0.055/—)	0.0131 (0.002/0.002/—)	0.0020 (0.200/0.341/—)
DDR	0.0445 (<0.001/<0.001/—)	0.0350 (<0.001/<0.001/—)	0.0146 (0.018/0.034/—)
Circuit × year + issue FE	✓	✓	✓
Panel controls	✓	✓	✓
Observations	96,849	96,849	96,849
R ²	0.0189	0.0124	0.0089

Table A.11: Main text and permutation p -values for between-panel model of Trump appointees’ effects on separate writing, following the between-panel model of Table 5.

	Dissent	Concurrence	Mixed
<i>Trump-R – non-Trump-R</i>			
RRR	-0.0005 (0.853/0.873/0.874)	0.0068 (0.079/0.102/0.076)	0.0028 (0.063/0.065/0.061)
RRD	0.0044 (0.263/0.406/0.799)	0.0141 (<0.001/<0.001/<0.001)	0.0041 (0.020/0.035/0.008)
DDR	0.0470 (<0.001/<0.001/—)	0.0375 (<0.001/<0.001/—)	0.0133 (0.011/0.046/—)
Case FE	✓	✓	✓
Observations	96,849	96,849	96,849
R ²	0.3259	0.3224	0.3305

Table A.12: Main text and permutation p -values for within-panel model of Trump appointees’ effects on separate writing, following the within-panel model of Table 5.

A.4 Extensions

A.4.1 Salient Topics

Our focus so far has been on the universe of published cases heard by the Courts of Appeals. Given that the academic research on Trump’s judges has emphasized particularly salient topics, in this section we consider three high-salience domains—gun rights (Brown, Epstein and Gulati 2025), religious liberty (Rothschild 2022, Choi, Gulati and Posner 2025), and immigration. For each, we select relevant cases using Songer and SCDB issue area sub-codes from the CCAD, administrative codes from the Federal Judicial Center’s Integrated Database (Federal Judicial Center 2026), and a simple regular-expressions (regex) search in the opinion text. For gun rights cases, we further separate into civil challenges to gun restrictions and gun-related claims by criminal defendants using the CCAD’s broad issue area code for Criminal cases. For each issue area, we report three outcome measures: the same ideological direction coding used throughout the paper, a CCAD-derived proxy using focal-party-related fields to reconstruct whether the claimant prevailed in each case, and an LLM-coded claimant-won variable.⁵ Details of the regex search and LLM coding pass are included following our discussion of results.

We report three comparisons: the raw difference in means, the coefficient on Trump appointees from a between-case regression with circuit-by-year fixed effects and panel controls, and the coefficient from a within-case regression with case fixed effects (absorbing and extending all previous controls). All three are judge-level comparisons, following the main text (Table 4) in assigning each judge on a panel the majority outcome when they support it and the opposite when they dissent. In each, we show a single Trump effect pooled across all panels because the limited sample size makes estimating interaction terms with the total number of Republican appointees (as in the universe-level Table 4) infeasible. The raw dif-

⁵For these topics, direction often disagrees with the other two measures. For example, according to the Songer codebook a rights claimant prevailing is coded as liberal even if they are invoking their Second Amendment right to bear arms.

ference in means is closest to the prior literature, which often includes some adjustments for demographic differences between Trump appointees and other judges. However, only the two fixed-effects models are causally identified, though their statistical power is limited by sample size. The number of judge-votes is shown in the “N (ID)” row; the raw differences include only Republican-appointed judges and thus have a lower count, while the parenthetical for the case fixed effects model shows the number of cases with identifying variation produced by differences among judges on the appropriate outcome variable. While the circuit-by-year fixed effects model generally has enough observations to produce precision comparable to the main text’s universe-level (and sharper) case fixed effects model, the case fixed effects model often relies on fewer than 100 cases for identification.

In Table A.13, we first compare Trump and non-Trump Republicans on the ideological direction and proxy measures (we return to the LLM measure shortly). For religious liberty cases, we use another regex search to separate cases which include Christian-related words, Muslim-related words, both, or neither. The raw differences in Table A.13 are consistent with prior literature—as captured by our CCAD-derived proxy, Trump appointees are more pro-gun, more pro-Christian, less pro-Muslim, and less pro-immigrant. Point estimates from the two causally identified regression analyses mostly support these results, but are generally not statistically significant. The main exception is immigration, which has nontrivial identifying variation even in the case fixed effects model; that model finds statistically significant results at the 0.1% level on both outcomes.

In Table A.14, we show results using the same approaches but narrowing the sample to the subset of cases where the LLM identified a genuine claim and replacing the CCAD-derived proxy with the LLM’s claimant-won code. The raw means again generally match prior literature—Trump appointees are more pro-gun, more pro-Christian (though also more pro-Muslim in this sample), and less pro-immigrant. The two causally identified models deliver stronger results in this table than in the previous one, though with the caveat that the smaller sample makes identifying variation in the case fixed effects model suspect. All

	Raw diff		Circuit $\times$ year FE		Case FE	
	Proxy	Dir.	Proxy	Dir.	Proxy	Dir.
Gun rights (all)	+1.73	-0.47	+0.29	+0.61	+0.30	+0.67
	—	—	(1.32)	(1.34)	(0.72)	(0.79)
N (ID)	4,656	4,656	7,557	7,557	7,557 (222)	7,557 (256)
civil 2A challenge	+3.90	+4.87	-0.60	+8.82*	+4.03	+5.08
	—	—	(3.92)	(4.33)	(2.90)	(3.61)
N (ID)	524	524	957	957	957 (57)	957 (78)
criminal 2A defense	+0.37	-0.24	+0.51	-0.82	-0.31	+0.03
	—	—	(1.20)	(1.24)	(0.73)	(0.74)
N (ID)	4,132	4,132	6,600	6,600	6,600 (165)	6,600 (178)
Religious liberty	+4.75	-4.26	+1.96	-1.42	+2.50	-1.55
	—	—	(2.50)	(2.45)	(2.04)	(2.28)
N (ID)	1,229	1,229	2,179	2,179	2,179 (100)	2,179 (135)
Christian (regex)	+3.16	-4.68	+1.86	-1.34	+2.70	-2.15
	—	—	(3.59)	(3.12)	(3.06)	(3.28)
N (ID)	625	625	1,111	1,111	1,111 (59)	1,111 (76)
Muslim (regex)	-9.29	+9.29	-16.91*	+12.85	-6.15	+1.62
	—	—	(7.95)	(9.66)	(6.17)	(7.74)
N (ID)	56	56	105	105	105 (2)	105 (3)
Immigration	-2.11	+1.98	-3.32	+2.48	-3.29***	+3.13***
	—	—	(1.75)	(1.79)	(0.80)	(0.90)
N (ID)	5,170	5,170	9,541	9,541	9,541 (313)	9,541 (386)

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

Standard errors in parentheses, two-way clustered by case and judge. “Circuit $\times$ year FE” columns include panel controls (indicators for number of Democratic appointees on the panel, ≥ 1 woman on the panel, and ≥ 1 nonwhite judge on the panel, as well as the panel’s mean age); those controls are absorbed by the case fixed effects in “Case FE” columns.

The sample includes all cases identified by Songer or SCDB issue area subcode, IDB code, or regex search. Civil and criminal gun rights cases are separated using the Songer broad issue area code for Criminal from the CCAD; Christian and Muslim religious liberty plaintiffs are separated using regex search.

Table A.13: Comparison of Trump appointees and non-Trump-Republican appointees on salient topics using existing CCAD data. The three column blocks provide the descriptive difference in means; the coefficient from a between-case regression with circuit-by-year fixed effects and panel controls; and the coefficient from a within-case regression with case fixed effects.

point estimates on the claimant-won code match prior literature (now including a negative coefficient for Muslim religious liberty claimants) and most are statistically significant in the case fixed effects model.

We draw three broad conclusions from these results. First, and most importantly, while statistical power is limited in this setting, causally identified estimates do find differences between Trump appointees and other Republican appointees on salient topics. Point esti-

	Raw diff		Circuit×year FE		Case FE	
	Claim	Dir.	Claim	Dir.	Claim	Dir.
Gun rights (all)	+10.00	-5.94	+6.65*	-2.21	+5.51*	+0.33
	—	—	(3.18)	(3.25)	(2.49)	(2.27)
<i>N</i> (ID)	971	971	1,585	1,585	1,585 (57)	1,585 (63)
civil 2A challenge	+22.41	-6.14	+8.82	+5.42	+11.81*	+3.70
	—	—	(6.58)	(6.71)	(5.85)	(5.89)
<i>N</i> (ID)	213	213	352	352	352 (26)	352 (29)
criminal 2A defense	+5.74	-5.28	+6.49*	-4.53	+3.53	-0.50
	—	—	(2.85)	(3.22)	(2.37)	(2.11)
<i>N</i> (ID)	758	758	1,233	1,233	1,233 (31)	1,233 (34)
Religious liberty (all)	+7.59	-6.48	+2.98	-1.17	+4.35*	-1.40
	—	—	(2.48)	(2.82)	(2.02)	(2.21)
<i>N</i> (ID)	727	727	1,294	1,294	1,294 (51)	1,294 (56)
Christian claimant	+7.92	-7.43	+3.55	+2.72	+8.06*	-2.72
	—	—	(4.07)	(3.85)	(3.26)	(3.61)
<i>N</i> (ID)	343	343	583	583	583 (30)	583 (33)
Muslim claimant	+2.69	-6.23	-9.11	+8.24	-3.28	+1.49
	—	—	(8.22)	(7.33)	(6.94)	(6.77)
<i>N</i> (ID)	98	98	194	194	194 (8)	194 (8)
Immigration	-1.43	+1.93	-1.04	+1.87	-1.91*	+2.73**
	—	—	(2.03)	(2.00)	(0.83)	(0.92)
<i>N</i> (ID)	3,196	3,196	5,762	5,762	5,762 (152)	5,762 (193)

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

Standard errors in parentheses, two-way clustered by case and judge. “Circuit × year FE” columns include panel controls (indicators for number of Democratic appointees on the panel, ≥ 1 woman on the panel, and ≥ 1 nonwhite judge on the panel, as well as the panel’s mean age); those controls are absorbed by the case fixed effects in “Case FE” columns.

The sample is the subset of cases from Table A.13 which the LLM flagged as containing the relevant claim. Civil and criminal gun rights cases, as well as Christian and Muslim plaintiffs, are separated by LLM codings.

Table A.14: Comparison of Trump appointees and non-Trump-Republican appointees on salient topics using novel LLM codings (CCAD-derived direction included for comparison); the structure follows Table A.13.

mates match the raw differences presented in prior literature almost across the board, and the sharpest test—a case fixed effects model using LLM codings to drop non-relevant cases from our sample—finds statistically significant effects for civil gun rights cases, Christian religious liberty claimants, and immigration cases. Second, ideological direction does not cleanly capture the outcome of interest; because it matches the claimant-won measure on some cases but not on others, it can alter the magnitude of effects and even (for religious liberty and immigration cases) flip their sign. This disagreement is a known weakness of

the Songer codebook’s approach to ideological direction, and should be taken into account when working with particular subsets of cases. Finally, the CCAD closely tracks a tailored LLM pass, showing the strength and flexibility of its comprehensive data. The main gap is over-including cases where the key issue is mentioned but not actually central, which can likely be fixed with a narrower opinion-text search.

A.4.1.1 Methods in Detail

All regex terms allow case-insensitive matches and are presented in a semicolon-separated list under a header describing which issue area they are designed for. For analysis of Christian and Muslim claimants in Tables A.13 and A.14, we require that only words from the appropriate list be present (i.e., no Christian regex matches can appear in a case designated as having a Muslim claimant, and vice-versa).

GUN-RIGHTS OPINION-TEXT DETECTOR

second amendment ; right to (:keep and)?bear arms ; \bbruen\b ; district of columbia
↪ v\.? heller ; mcdonald v\.? city

RELIGIOUS-LIBERTY OPINION-TEXT DETECTOR

free exercise ; establishment clause ; \bRFRa\b ; \bRLUIPA\b ; religious freedom
↪ restoration ; religious liberty ; freedom of religion ; religious exercise

CHRISTIAN CLAIMANT DETECTOR

\bchristian(?:s|ity)?\b ; \bchrist\b ; \bchurch(?:es)?\b ; \bcatholic\b ; \bbaptist\b ;
↪ \bevangelical\b ; \bprotestant\b ; \bcongregation(?:al)?\b ; \bparish\b ; \bdiocese\b
↪ ; \bpastor\b ; \bgospel\b ; \bbible\b ; \bbiblical\b ; \bjesus\b ; \bmethodist\b ;
↪ \bpresbyterian\b ; \blutheran\b ; \bpentecostal\b ; \bmennonite\b ; \bamish\b ;
↪ seventh-day adventist ; \bcrucifix\b

MUSLIM CLAIMANT DETECTOR

\bmuslim(?:s)?\b ; \bislam(?:ic|ist)?\b ; \bmosque\b ; \bqur'?an(?:ic)?\b ; \bkoran\b ;
↪ \bramadan\b ; \bhalal\b ; \bimam\b ; \ballah\b ; \bhijab\b ; \bsharia\b ; \bsunni\b ;
↪ \bshia\b ; prayer rug

IMMIGRATION OPINION-TEXT DETECTOR

board of immigration appeals ; immigration judge ; immigration and nationality act ;

↪ removal proceeding ; removal order ; withholding of removal ; cancellation of removal

↪ ; notice to appear ; \basylum\b ; convention against torture ; \bdeportation\b

Claim coding for gun rights, religious liberty, and immigration cases is produced by Qwen3-235B-A22B-Instruct-2507 (Yang et al. 2025), the same open-weights model used by the CCAD. The model has 235 billion total parameters (22 billion activated per token) and a 128,000-token context window, large enough to hold all but a handful (fewer than 100) of opinions in our sample. Any long opinions are head-and-tail truncated, using 75% of the token budget for the beginning and 25% for the end to ensure both introductory and concluding material is preserved. We use the instruction-tuned variant with no fine tuning and set temperature to zero for reproducibility. Both tasks run through [DeepInfra's](#) OpenAI-compatible API. Each query is independent with no system memory.

We ask the model to read a single opinion that our filter (the combination of issue area codes and regex described at the beginning of the section) has flagged as possibly involving the relevant claim, and to code—from the perspective of the party asserting the right (the “claimant”)—whether the claim is actually adjudicated on the merits, what type of claim it is, who the claimant is, and whether the claimant prevailed. A shared system prompt is specialized to each domain; the model returns a JSON object with the requested variables. The gun rights prompt is reproduced below.

You are an expert analyst of United States Court of Appeals opinions.

You will be given the text of ONE opinion. The case has been flagged as possibly involving a Second Amendment / right-to-keep-and-bear-arms claim. Read the opinion and
↪ code how it resolves that claim,

FROM THE PERSPECTIVE OF THE PARTY ASSERTING THE RIGHT (the "claimant").

Code these fields:

- claim_present: true if the opinion actually adjudicates a Second Amendment /
↳ right-to-keep-and-bear-arms claim on the merits (not merely a passing mention, citation, or unrelated use); false otherwise. If false, set the other fields to "na"/empty.
- claim_type: one of
"civil_challenge" = a civil action challenging a firearms law, regulation, or
↳ permit/license denial (the claimant affirmatively seeks to vindicate gun rights)
"criminal_defense" = a criminal defendant invoking the Second Amendment to contest a
↳ firearms charge, conviction, or sentence
"other" = the Second Amendment is mentioned but the case is neither of the above
- claimant: which side is the party asserting the right --- "appellant" or "respondent" (the appellant/petitioner is the party who brought the appeal).
- claimant_desc: a SHORT (<= 8 words) description of that claimant (e.g., "convicted felon", "gun-rights association", "Muslim prisoner", "asylum seeker from Honduras").
- claimant_won: did the panel rule IN FAVOR OF the claimant's claim?
"yes" | "no" | "mixed" (partial) | "na" (no merits ruling).

GUARDRAILS

- Use only the opinion text. Do NOT speculate about the judges' identities, party, or who appointed them.
- "claimant" identifies who asserts the right, NOT who won.
- If the opinion text is truncated, code from what is provided.

OUTPUT

Output ONLY a single JSON object, no prose before or after, with EXACTLY these keys:

```
{
  "claim_present": true,
  "claim_type": "...",
  "claimant": "appellant" | "respondent",
  "claimant_desc": "...",
  "claimant_won": "yes" | "no" | "mixed" | "na"
}
```

The religious liberty and immigration prompts are identical except for the substituted claim description, claim-type menu, and (for religion) the additional `claimant_religion` field:

RELIGION

claim: a Free Exercise Clause, Establishment Clause, RFRA, or RLUIPA claim

claim_type:

"free_exercise" = a free-exercise (or RFRA/RLUIPA) claim that a government action burdens

↪ religious exercise

"establishment" = an Establishment Clause claim (government endorsement or support of

↪ religion)

"both" = both free-exercise and establishment claims are at issue

"other" = religion is mentioned but the case is neither of the above

extra field claimant_religion: the religion of the party whose religious exercise or

↪ interest is at issue: "christian" | "jewish" | "muslim" | "other" | "multiple" |

↪ "unclear"

IMMIGRATION

claim: a noncitizen's claim for immigration relief or against removal

claim_type:

"asylum" = asylum

"withholding" = withholding of removal

"CAT" = protection under the Convention Against Torture

"cancellation" = cancellation of removal

"adjustment" = adjustment of status, visa, or admissibility

"removability" = whether the noncitizen is removable (a charge of

↪ removability/deportability)

"other" = an immigration matter not covered above

A.4.2 Other Presidents' Appointees

Our design isolates whether Trump appointees vote differently from their co-partisan appointees; a natural extension is to apply this design to other presidents. We re-estimate

case outcome (Table 3) and judge vote (Table 4) models with the focal cohort changed from Trump appointees to (i) George W. Bush appointees or (ii) Biden appointees. The Bush comparison asks whether another recent Republican appointing-president cohort shows a similar within-party pattern; because Bush began appointing judges in 2001, we extend the sample window to 2001–2025, which also yields far more identifying variation than the post-2017 Trump sample. The Biden comparison is intended as a symmetric test on the Democratic side and uses only data from 2021–2025 since no Biden appointees exist prior to 2021.

Tables A.15 and A.16 report the Bush results. The Bush-appointee coefficient on case outcomes is never statistically significant and point estimates never exceed one percentage point. At the same time, partisan composition effects remain large and highly significant—an all-Republican panel is roughly 16 percentage points more likely to reach a conservative outcome than an all-Democratic one. Point-estimated differences between Bush appointees and other Republican appointees are about the same size as for Trump appointees on RRD panels, but much smaller on DDR panels, where the main text found its largest effect. The within-case design has about four times as many identifying panels as the Trump appointee within-case design (over 1,600 panels where Bush and non-Bush-Republican appointees vote differently), making the implied minimum detectable effect a fraction of a percentage point.

Tables A.17 and A.18 report the Biden results. The Biden-appointee case outcome effect is not statistically significant, and the partisan difference—DDD panels produce about 22 percentage points fewer conservative outcomes than RRR panels—is an order of magnitude larger than any point estimate. The between-panel point estimates for DDR and RRD panels are comparable to many of the Trump-appointee effects in the main text—about 1–2 percentage points—but are not statistically significant, while the within-panel estimates for all panel types are all smaller than one percentage point and again not statistically significant. The reduced sample size and thin within-case identification—fewer than 100 cases feature a Biden appointee and a non-Biden-Democratic appointee voting differently—means that all coefficients carry relatively large standard errors (generally about 2 percentage points).

While neither of these extensions shows a distinctive appointing-president cohort, it would have been unexpected to find large differences for either of these two presidents. However, the results themselves may also be of independent interest, especially given the key role Bush appointees have played in the current Supreme Court's conservative supermajority and the Democratic Party's renewed emphasis on judicial nominations at the district and circuit level given their perceived defeats at the Supreme Court.

	All panels		Circuit judges only	
	(1)	(2)	(3)	(4)
#Bush-R seat $\times$ RRR	0.0036 (0.0045)	-0.0024 (0.0046)	0.0030 (0.0049)	-0.0040 (0.0051)
#Bush-R seat $\times$ RRD	0.0077 (0.0040)	0.0035 (0.0041)	0.0054 (0.0045)	0.0008 (0.0045)
#Bush-R seat $\times$ DDR	-0.0015 (0.0074)	-0.0067 (0.0071)	-0.0034 (0.0081)	-0.0086 (0.0078)
RRR panel	0.1592*** (0.0100)	0.1617*** (0.0107)	0.1591*** (0.0111)	0.1602*** (0.0116)
RRD panel	0.1259*** (0.0080)	0.1290*** (0.0083)	0.1259*** (0.0092)	0.1281*** (0.0093)
DDR panel	0.0653*** (0.0070)	0.0695*** (0.0072)	0.0627*** (0.0081)	0.0668*** (0.0081)
Circuit $\times$ year FE	✓	✓	✓	✓
Broad issue area FE		✓		✓
Panel controls		✓		✓
Observations	105,881	105,881	88,057	88,057
R ²	0.0414	0.0848	0.0427	0.0879

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

Table A.15: Differences in conservative case outcomes by George W. Bush appointees vs. other Republican appointees, 2001–2025 (following Table 3).

	Between panels		Within panels	
	All panels (1)	Circuit judges only (2)	All panels (3)	Circuit judges only (4)
<i>Bush-R – non-Bush-R</i>				
RRR	-0.0018 (0.0068)	-0.0027 (0.0071)	0.0002 (0.0019)	-0.0000 (0.0020)
RRD	0.0080 (0.0068)	0.0070 (0.0071)	0.0042 (0.0029)	0.0043 (0.0031)
DDR	-0.0018 (0.0099)	-0.0033 (0.0104)	0.0074 (0.0060)	0.0084 (0.0064)
<i>non-Bush-R rate (modeled)</i>				
RRR	0.6388	0.6411	0.6679	0.6698
RRD	0.6148	0.6181	0.6186	0.6206
DDR	0.5816	0.5876	0.5591	0.5617
Case FE			✓	✓
Circuit $\times$ year + issue FE	✓	✓		
Panel controls	✓	✓		
Observations	317,643	299,732	317,643	299,732
R ²	0.0799	0.0808	0.9181	0.9196

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

Table A.16: Judge-level differences in conservative votes focusing on George W. Bush appointees, 2001–2025 (following Table 4).

	All panels		Circuit judges only	
	(1)	(2)	(3)	(4)
#Biden-D seat $\times$ DDD	0.0091 (0.0153)	0.0075 (0.0168)	0.0152 (0.0163)	0.0184 (0.0173)
#Biden-D seat $\times$ DDR	-0.0143 (0.0171)	-0.0160 (0.0190)	-0.0146 (0.0180)	-0.0137 (0.0198)
#Biden-D seat $\times$ RRD	0.0090 (0.0214)	0.0101 (0.0208)	0.0099 (0.0219)	0.0139 (0.0212)
DDD panel	-0.2180*** (0.0249)	-0.2221*** (0.0257)	-0.2182*** (0.0277)	-0.2236*** (0.0291)
DDR panel	-0.1480*** (0.0168)	-0.1433*** (0.0166)	-0.1527*** (0.0185)	-0.1485*** (0.0183)
RRD panel	-0.0523*** (0.0138)	-0.0525*** (0.0130)	-0.0535*** (0.0143)	-0.0542*** (0.0135)
Circuit $\times$ year FE	✓	✓	✓	✓
Broad issue area FE		✓		✓
Panel controls		✓		✓
Observations	17,176	17,176	15,760	15,760
R ²	0.0467	0.1003	0.0472	0.1038

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

Table A.17: Differences in conservative outcomes by Biden appointees vs. other Democratic appointees, 2021–2025 (following Table 3).

	Between panels		Within panels	
	All panels (1)	Circuit judges only (2)	All panels (3)	Circuit judges only (4)
<i>Biden-D – non-Biden-D</i>				
DDD	0.0053 (0.0206)	0.0027 (0.0209)	-0.0026 (0.0049)	-0.0046 (0.0051)
DDR	-0.0182 (0.0175)	-0.0177 (0.0177)	0.0032 (0.0060)	0.0038 (0.0062)
RRD	0.0165 (0.0227)	0.0171 (0.0228)	0.0062 (0.0155)	0.0067 (0.0158)
<i>non-Biden-D rate (modeled)</i>				
DDD	0.4624	0.4664	0.4895	0.4939
DDR	0.5229	0.5214	0.5055	0.5053
RRD	0.5943	0.5926	0.5880	0.5888
Case FE			✓	✓
Circuit $\times$ year + issue FE	✓	✓		
Panel controls	✓	✓		
Observations	51,528	50,112	51,528	50,112
R ²	0.0947	0.0960	0.9149	0.9162

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

Table A.18: Judge-level differences in conservative votes focusing on Biden appointees, 2021–2025 (following Table 4)

Appendix References

- Brown, Rebecca L., Lee Epstein and Mitu Gulati. 2025. “Guns, Judges, and Trump.” *Duke Law Journal Online* 74:81–109.
- Choi, Stephen J., Mitu Gulati and Eric A. Posner. 2025. “Trump’s Lower Court Judges and Religion: An Initial Appraisal.” *The Journal of Legal Studies* 54(1):1–41.
- Chung, EunYi and Joseph P. Romano. 2013. “Exact and Asymptotically Robust Permutation Tests.” *The Annals of Statistics* 41(2):484–507.
- Federal Judicial Center. 2026. “Federal Court Cases: Integrated Database.” <https://www.fjc.gov/research/idb>. Data updated regularly; exported data current through 31 December 2025.
- Rothschild, Zalman. 2022. “Free Exercise Partisanship.” *Cornell Law Review* 107:1067–1136.
- Sapiro-Gheiler, Eitan and Jonathan P. Kastellec. 2026. “Appealing to Large-Language-Model-as-Judge: A Comprehensive Machine-Coded Database for the U.S. Courts of Appeals.” Working paper, available at coadata.org.
- Yang, An, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, Chujie Zheng, Dayiheng Liu, Fan Zhou, Fei Huang, Feng Hu, Hao Ge, Haoran Wei, Huan Lin, Jialong Tang, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiayi Yang, Jing Zhou, Jingren Zhou, Junyang Lin, Kai Dang, Keqin Bao, Kexin Yang, Le Yu, Lianghao Deng, Mei Li, Mingfeng Xue, Mingze Li, Pei Zhang, Peng Wang, Qin Zhu, Rui Men, Ruize Gao, Shixuan Liu, Shuang Luo, Tianhao Li, Tianyi Tang, Wenbiao Yin, Xingzhang Ren, Xinyu Wang, Xinyu Zhang, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yinger Zhang, Yu Wan, Yuqiong Liu, Zekun Wang, Zeyu Cui, Zhenru Zhang, Zhipeng Zhou and Zihan Qiu. 2025. “Qwen3 Technical Report.”. arXiv preprint, <https://arxiv.org/abs/2505.09388>. Model access: <https://huggingface.co/Qwen/Qwen3-235B-A22B-Instruct-2507>.