

# Politics, Elections, and Coercive Bureaucracy\*

Tom S. Clark<sup>†</sup>

August 31, 2026

## Abstract

Political control of bureaucracies runs both through hierarchical leadership mechanisms and through environmental factors that shape the conditions under which decisions are made. While rich theoretical arguments illuminate these sources of change and the mechanisms through which they operate, they can be difficult or impossible to empirically distinguish. I offer novel insight into the distinction between electoral and environmental sources of change and the institutional mechanisms through which they operate. Using comprehensive data from Cook County linking policing and prosecution records, I study the impact of the election of a reform prosecutor on how discretion is exercised in the context of felony crimes. I show that the election of a new leader had only a minimal impact on felony prosecution, limited to narrow formal policies adopted, whereas broader environmental forces shaped discretionary decision-making independent of the formal leadership change. These results show how political control can operate through incumbents and begin before formal succession and have implications for empirical studies focusing on the impact of electoral turnover.

**Preliminary and subject to errors. Please do not circulate or cite.**

---

\*I thank Jon Rogowski for helpful comments.

<sup>†</sup>Professor of Political Science, Stanford University

# 1 Introduction

Prosecutors exercise tremendous power over American life, wielding nearly unfettered discretion over the application of criminal law. In the US, most prosecutors are elected ([Lerman and Weaver 2014](#), [Weaver and Lerman 2010](#)), and those who are not nevertheless typically are directly accountable to elected politicians. Over the past decade, voters across the country have elected candidates promising to change how that power is exercised. And, in the federal system, we have seen increasing concern about federal prosecutors' independence from political influence. Yet prosecutors sit atop bureaucratic organizations in which line attorneys exercise considerable discretion over individual cases. This creates a familiar problem of political control: how and under what conditions do political leaders shape the conduct of the bureaucrats who implement their policies?

Past research has taught us a great deal about the tools elected officials use to hold bureaucracies in check and about when those tools bind ([Gailmard and Patty 2012](#), [McCubbins, Noll and Weingast 1987](#), [McCubbins and Schwartz 1984](#), [Wood and Waterman 1991](#)). However, we know less about what patterns of change in a bureaucracy reveal about those various mechanisms or how bureaucratic change can distinguish changes in leadership from changes in the broader political environment. In particular, a transition in leadership may be politically consequential without itself being the root cause or the temporal point of bureaucratic change.

Among political scientists, the dominant account identifies political leaders as the source of bureaucratic change. The principal-agent tradition highlights the oversight powers a leader has as a mechanism that can induce agents to change how they exercise discretion. A new leader can replace personnel or announce rules. She can also move the incumbents who remain through supervision and the prospect of reassignment ([Bolton, de Figueiredo and Lewis 2021](#), [Sasso 2020](#), [Weingast and Moran 1983](#), [Wood and Waterman 1991](#)). An alternative account points to the environment in which bureaucrats exercise discretion. A scandal or burst of public scrutiny can change what an office can afford to be seen doing,

even before a new leader arrives (Bustos 2021, Busuioc and Lodge 2016, Carpenter 2010). From this perspective, bureaucratic change may instead reflect a shift in the environment in which officials exercise discretion. Thus, an analysis that focuses on turnover in political leadership can misattribute change to a new leader when in fact it was part of changes in the political environment surrounding the transition. These processes are easy to conflate because the forces that reshape an agency’s environment often change who leads it. Career officials may also anticipate the principal they expect to serve. An analysis that focuses on the impact of political turnover is therefore liable to misattribute change brought about with the surrounding political environment centered to change in leadership.

Prosecutors are an unusually revealing setting in which to study this problem. They decide whom to charge and on what terms, their decisions are largely unreviewable, and in nearly every state they answer directly to voters rather than to an administrative superior (Gordon and Huber 2002).<sup>1</sup> Recent studies use reform prosecutors taking office to estimate changes in charging and conviction (Agan et al. 2025) or imprisonment (Gottlieb et al. 2025). Others examine crime (Petersen, Mitchell and Yan 2024) or isolate the effects of marginal charging decisions (Agan, Doleac and Harvey 2023). These studies illuminate important consequences of prosecutorial reform. However, changes in criminal justice outcomes following electoral turnover do not necessarily isolate the effect of the elected prosecutor. Choices made by police officers, for example, determine which cases come to prosecutors in the first instance, so a charging rate reflects both what prosecutors decide and the cases police bring forward (Clark 2026a, Knox, Lowe and Mummolo 2020). The office may also have begun adapting before the prosecutor takes office.

I offer a simple framework that separates the source of changes in prosecution from the institutional mechanism through which that change occurs. In particular, changes in how prosecutors handle cases may stem from either the preferences of a current or anticipated

---

<sup>1</sup>The United States is nearly alone in electing prosecutors, a practice used in 45 states. Local prosecutors are appointed at the state level in Alaska, Connecticut, and New Jersey; in Delaware and Rhode Island, prosecutions are handled centrally by the elected Attorney General.

prosecutor or changes the office’s reputational and political environment. The framework contemplates three mechanisms. The first is personnel, which involves a leader making choices about which line attorneys handle cases. Personnel change alters the distribution of stable reviewer types, whereas movement in the common standard appears within reviewers over their observed tenures. A second is policy, focusing on how new directives or rules can constrain discretion. Rules about which cases ay or may not be prosecuted should produce a sharp response among covered cases. The third is movement in the common standard reviewers apply, either due to direction from a leader or a response to the office’s reputational environment. Moreover, if a change in political leadership initiates the office-wide movement, approval changes level or slope at the primary or inauguration, whereas anticipation implies an earlier beginning to the change. Distinguishing these mechanisms reveals how political control operates: through the selection of officials, formal restrictions on their choices, or changes in how incumbents exercise the discretion they retain.

I evaluate these implications in the Cook County State’s Attorney’s Office from 2014 through 2019. Kim Foxx defeated the incumbent in the March 2016 Democratic primary and took office that December after campaigning to reform the office. I focus on felony review, which is the first decision point and involves on-call assistant state’s attorney approving or rejecting the felony charges police present. Incoming cases rotate among active reviewers in a manner close to random conditional on the roster and call time ([Jordan 2024](#)). That rotation lets me separate changes in the measured distribution of reviewers from changes in how reviewers decide. I also employ a formal policy Foxx adopted that the office would decline to prosecute retail theft under \$1000 as a felony as a benchmark for formal policy. I also link Chicago arrests to review and charging records. This allows me to distinguish the cases police present to prosecutors from the decisions prosecutors make.

Over the period covered here, prosecutors declined in their propensity to approve felony charges, and most of that decline appears as common change within reviewers rather than a shift in the mix of reviewers. The office’s retail-theft rule produces the sharp, offense-specific

response predicted by a categorical policy, but it does not explain the broader path. That change was already under way before Foxx took office: neither the Democratic primary nor her inauguration is associated with its onset, although the estimates do not rule out gradual influence after office-taking. Changes in the case pool also do not account for the decline. The decline therefore appears in the common standard reviewers applied rather than in the mix of reviewers or cases. Foxx’s inauguration is therefore a poor treatment date for identifying her contribution. More broadly, the results show that bureaucratic change can occur through incumbent adaptation before a transition in leadership. Accounts of political control that begin with office-taking may therefore begin too late.

## 2 How Political Change Reaches a Bureaucracy

Like virtually every bureaucracy in government, political control of a prosecutor’s office operates through delegation. Voters (in elective systems) select a prosecutor who in turn delegates to staff attorneys the responsibility for handling individual cases ([Gordon and Huber 2002](#)).<sup>2</sup> In every jurisdiction save for the smallest ones, there are simply too many cases for the elected prosecutor to handle herself. She instead shapes how line attorneys exercise discretion through personnel authority and supervision, backed by the prospect of reassignment ([Gailmard and Patty 2012](#), [McCubbins, Noll and Weingast 1987](#), [Weingast and Moran 1983](#), [Wood and Waterman 1991](#)). These tools can change who makes decisions, but they can also affect how attorneys use their discretion in the first instance. In the spirit of canonical principal-agent models, attorneys who expect closer supervision or reassignment have strong reasons to adapt, even though their own sincere preferences remain constant. Political control can therefore appear as change within incumbent officials rather than as personnel replacement.

The possibility of within-attorney adjustment means that a political transition a process

---

<sup>2</sup>There are, of course many exceptions to this general rule. Most notably, the chief prosecutor, typically called a district attorney, may be especially involved in some politically salient cases.

rather than a single intervention. Bureaucrats respond to the principal who currently oversees them, but they may also respond to the principal they expect to serve. As an election season unfolds, career officials learn about the likely successor and the preferences she is expected to bring. Anticipation can therefore change conduct before the new principal possesses formal authority (Bolton, de Figueiredo and Lewis 2021, Doherty, Lewis and Limbocker 2019, Sasso 2020). Of course, the incumbent may respond as well, facing an electoral threat that creates an incentive to change how the office operates before voters decide whether she will retain control. Further, even after an election that elads to a change in leadership, the new leader’s impact may take a while to be felt as supervision and training alter how line officials understand the priorities of the office (Oberfield 2014). A new leader may therefore come into office in the midst of an institutional response already under way. Analyses focusing on the effects of turnover in an office may misstate the timing of political influence even when the incoming leader eventually matters.

But anticipation of turnover is only one reason an office may change before its leadership does. A bureaucracy may very well be sensitive to perceptions among the audiences that judge it. For example, scandals may increase the reputational cost of decisions perceived to be wrong or illegitimate, changing what officials regard as defensible without a directive from above (Birkland 1998, Carpenter 2010, Gilad, Maor and Ben-Nun Bloom 2015, Kingdon 1984). In the context of a prosecutor’s office, that means line attorneys may exercise the same formal authority differently because the environment in which they exercise it has changed. When the leader is an elected official, the same event may also discredit an incumbent and alter the prospects of a challenger. Thus, political events may change both the office’s reputational environment and its expectations about future leadership. Changes we observe in bureaucratic decision-making can therefore reflect either pressure—a distinction that timing alone does not make.

I distinguish these sources of change from the institutional mechanisms through which change occurs. Hierarchical pressure originates with the current or anticipated prosecutor.

Environmental pressure originates with the audiences evaluating the office.<sup>3</sup> The first mechanism is personnel composition: changing which attorneys exercise discretion (e.g., [Lewis 2008](#), [Wood and Waterman 1991](#)). The second is a categorical rule that constrains discretion for the cases it covers (e.g., [Huber and Shipan 2002](#), [McCubbins, Noll and Weingast 1987](#)). The third is the common standard reviewers apply (e.g., [Brehm and Gates 1997](#), [Oberfield 2014](#)). When the common standard changes, incumbent reviewers use their discretion differently, without a new rule or a change in personnel. That change may reflect direction from a leader or a response to the office’s political environment.

These mechanisms imply different forms of political control. If change occurs through personnel composition, then the prosecutor changes office behavior by changing which attorneys exercise discretion, rather than incumbent attorneys responding to supervision. If a prosecutor can bring about sharp changes via formal policy announcements, then perhaps she can control decisions in case she identifies *ex ante*, but that does not inform the exercise of discretion outside the policy. By contrast, changes that operate through the office’s common standard change how prosecutors handle cases without being replaced or having each decision prescribed by policy. Political control then operates through discretion that prosecutors exercise, due to effective supervision or anticipated sanctions, in the spirit of standard principal-agent models. However, this kind of change may also reflect reputational pressure, so it does not by itself identify the source of the common change.

### 3 A Dynamic Model of Felony Review

I build a simple dynamic model that connects political pressure to the charging decisions by line lawyers working within an office led by a prosecutor. The office persists across elections, while its elected principal (i.e., the prosecutor) may be retained or replaced. Reviewers (i.e., line lawyers) rotate through felony review and handle individual cases subject to a common organizational standard. The model separates changes in that standard from changes in the

---

<sup>3</sup>Of course, these sources are not exhaustive, but they are central to the political process I study here.

reviewers who staff the office.

**Players and characteristics.** The model has an office and a set of reviewers, indexed by  $a$ . Time is discrete,  $t = 1, 2, \dots$ , and cases are indexed by  $i$ . Case  $i$  has prosecutable strength  $v_i \in \mathbb{R}$ , an offense  $o_i$ , and characteristics  $x_i$ . Prosecutable strength includes the evidence, legal sufficiency, and other case-specific considerations relevant to felony charging. Conditional on  $x_i$ , it is drawn from a continuous distribution  $G(\cdot \mid x_i)$  known to the office and its reviewers.

In each period  $t$ , a set  $\mathcal{R}_t$  of reviewers is active, with membership that may change from one period to the next. Reviewer  $a$  has a preferred charging threshold  $\alpha_a \in \mathbb{R}$ , defined on the same scale as  $v_i$ . It is the minimum prosecutable strength she would require in the absence of organizational pressure. The threshold is fixed for each reviewer but varies across reviewers. Let  $\pi_{at}$  denote reviewer  $a$ 's share of cases in period  $t$ , with  $\pi_{at} = 0$  when  $a \notin \mathcal{R}_t$ . Let  $F_t(\alpha)$  denote the case-weighted distribution of preferred thresholds induced by these shares.

The elected prosecutor in office at  $t$  is  $q_t$ . Her intrinsic charging standard is  $\bar{s}_{q_t}$ . The identity of the principal changes only through an election, but the political state can change while she remains in office. Let  $Z_t$  contain the office's reputational and electoral state. The electoral state includes the incumbent's position, the calendar, and beliefs about future control. The state evolves according to a Markov transition that may depend on the standard the office implements. At an election, that transition determines whether  $q_t$  is retained or replaced. The office also inherits the equilibrium standard  $s_{t-1}^*$  implemented in the previous period. This state persists even when the principal changes and the reviewers rotate.

The incumbent's current ideal point is

$$\hat{s}_{q_t}(E_t) = \bar{s}_{q_t} + \phi_E E_t, \quad \phi_E > 0, \quad (1)$$

where  $E_t$  is the reputational component of  $Z_t$ . Reputational pressure may push the ideal point toward a more restrictive or more permissive standard. Electoral pressure enters

through the incumbent's continuation value: she values retaining office, and current policy may affect that probability. Finally,  $D_{it}$  equals one when a categorical rule is in force at  $t$  covers the case given its offense and any other eligibility conditions, including value or prior record.

**Choices and timing.** At the beginning of period  $t$ , the incumbent observes the inherited standard and the political state, then chooses an implemented common standard  $s_t \in \mathbb{R}$ . Each reviewer  $a \in \mathcal{R}_t$  observes that choice and chooses the charging threshold  $\tau_{at} \in \mathbb{R}$  she will apply. Cases then arrive and are assigned to active reviewers. The assigned reviewer observes  $v_i$  and approves the case when its prosecutable strength exceeds her chosen threshold. A binding categorical rule removes approval from the choice set when  $D_{it} = 1$ . The implemented standard becomes the organizational state inherited in period  $t + 1$ .

**Preferences.** I represent the reviewers' and office's preferences with loss functions. Reviewer  $a$  prefers a threshold close to her preferred charging threshold but also bears a cost from departing from the office standard. She chooses  $\tau_{at}$  to minimize

$$L_{at}(\tau_{at}; s_t) = (\tau_{at} - \alpha_a)^2 + \kappa (\tau_{at} - s_t)^2, \quad (2)$$

where  $\kappa \geq 0$  is the weight she places on the office standard relative to her own judgment. The parameter  $\kappa$  summarizes the career and supervisory consequences of departing from the office standard.

The incumbent is forward-looking. She discounts the future by  $\delta \in (0, 1)$  and bears the period loss

$$\ell_{qt}(s_t; s_{t-1}^*, E_t) = (s_t - \hat{s}_{qt}(E_t))^2 + \zeta (s_t - s_{t-1}^*)^2, \quad \zeta \geq 0, \quad (3)$$

where the second term makes rapid change costly.<sup>4</sup> Let  $V_q(s, Z)$  denote prosecutor  $q$ 's value

---

<sup>4</sup>Appendix A allows the office also to bear an implementation cost when the induced reviewer thresholds depart from reviewers' preferred standards. That extension changes the dynamic policy function but not the decomposition or the policy-scope implication evaluated below.

when the inherited standard is  $s$  and the political state is  $Z$ . Her Bellman equation is

$$V_q(s, Z) = \max_{s'} \left\{ -\ell_q(s'; s, E) + \delta \mathbb{E} \left[ \mathbf{1}\{q' = q\} V_q(s', Z') + \mathbf{1}\{q' \neq q\} \bar{V}_q \mid s, Z, s' \right] \right\}, \quad (4)$$

where  $\bar{V}_q$  is the incumbent's value after leaving office. Outside an election period,  $q' = q$  with probability one. At an election, the transition depends on the political state and may depend on the office's conduct. An incumbent can therefore change the standard before the election because doing so affects her prospect of remaining in office. A successor inherits the standard even though the departing incumbent does not value the successor's policy payoff.

**Equilibrium.** To solve the model, I characterize a Markov-perfect Stackelberg equilibrium. That means the incumbent's strategy depends on the inherited standard and the current political state, while reviewers observe the office's standards and best respond. [Appendix A](#) develops the dynamic program and gives a linear-quadratic benchmark.

Let  $w \equiv \kappa/(1+\kappa)$ . Given a candidate implemented standard  $s_t$ , reviewer  $a$ 's best response is

$$\tau_{at}(s_t) = (1 - w)\alpha_a + ws_t. \quad (5)$$

The incumbent's equilibrium policy solves Equation (4). I write it as

$$s_t^* = \sigma_{qt}(s_{t-1}^*, Z_t). \quad (6)$$

The policy function, in turn, balances the incumbent's current ideal point against the cost of changing the inherited standard, governed by  $\zeta$ . Its continuation term incorporates the value of retaining office and expectations about the political state that will follow. Substituting  $s_t^*$  into the reviewer best response gives the equilibrium threshold

$$\tau_{at}^* = (1 - w)\alpha_a + ws_t^*. \quad (7)$$

**Proposition 1** *Suppose the state and standard spaces are compact, payoffs are bounded and continuous, and the state transition is continuous. A Markov-perfect Stackelberg equilibrium exists. If the incumbent's Bellman objective has a unique maximizer, she uses the policy  $s_t^* = \sigma_{q_t}(s_{t-1}^*, Z_t)$ . Each reviewer  $a \in \mathcal{R}_t$  chooses  $\tau_{at}^* = (1 - w)\alpha_a + ws_t^*$ .*

In the linear-quadratic benchmark, the implemented standard partially adjusts toward the incumbent's ideal point within a fixed political regime. When office conduct affects retention, the incumbent may also adapt before an election in response to an electoral threat. The extension below allows the office standard to reflect reviewers' expectations about future control. Both mechanisms permit change before a formal transition in leadership. By contrast, an unexpected election outcome can produce an abrupt change in policy. Reviewer rotation changes the distribution of  $\alpha_a$  but does not alter the office's policy in the baseline model.

Allowing for a categorical rule, where  $D_{it} = 1$  when the rule applies to case  $i$  in period  $t$ , reviewer  $a$ 's equilibrium decision is therefore

$$Y_{ait} = \mathbf{1}\{v_i > \tau_{at}^*\} \mathbf{1}\{D_{it} = 0\}. \quad (8)$$

That is, the rule applies only to cases that satisfy its eligibility conditions.

**Composition and the common standard.** The model allows us to distinguish the effects of the composition of reviewers from the effect of a change in the common, office-wide standard. For cases outside a binding rule, a reviewer of type  $\alpha$  approves with probability

$$\mathcal{A}(\alpha, s_t^* \mid x_i) = 1 - G((1 - w)\alpha + ws_t^* \mid x_i). \quad (9)$$

Naturally, holding case characteristics fixed, approval becomes less likely when a stricter reviewer handles the case or when the office raises the common standard. Note that  $\mathcal{A}(\alpha, s_t^* \mid x_i)$  does not require a particular distribution for prosecutable strength. Although the threshold is linear in reviewer type and the common standard, the resulting approval probability

generally is not. A first-order approximation around the charging margin makes their contributions additive. Let  $\bar{f} > 0$  denote the density of prosecutable strength at that margin after standardizing the case mix. Then

$$\mathcal{A}(\alpha, s_t^*) \approx \mu - \bar{f}(1 - w)\alpha - \bar{f}ws_t^*. \quad (10)$$

Define the approval-scale reviewer effect  $\tilde{\alpha}_a = -\bar{f}(1 - w)\alpha_a$  and the common component  $g(t) = -\bar{f}ws_t^*$ . Office-wide approval is approximately  $A_t = \mu + \sum_{a \in \mathcal{R}_t} \pi_{at} \tilde{\alpha}_a + g(t)$ . Its change from a base period has the additive decomposition

$$\Delta A_t = \underbrace{\sum_{a \in \mathcal{R}_0 \cup \mathcal{R}_t} (\pi_{at} - \pi_{a0}) \tilde{\alpha}_a}_{C_t^F} + \underbrace{(g(t) - g(0))}_{C_t^S}. \quad (11)$$

The first term captures the change in the case-weighted mean reviewer effect. The second captures the change in the common component. Reviewer composition contributes to the personnel channel when case weights shift toward reviewers with different stable effects.

**Source-specific heterogeneity.** The baseline equilibrium characterizes a change in the common standard but does not identify which part of the political state produced it. Here I propose a modest extension to examine heterogeneity in how reviewers respond to hierarchical and political pressure. Specifically, this extension allows an additional response to vary with a predetermined measure  $z_a \in \mathbb{R}$  of reviewer  $a$ 's career exposure to that principal. Let  $B_t$  denote pressure from the current or anticipated principal. A common response to that pressure changes the office standard by  $\eta_B B_t$ , where  $\eta_B \geq 0$ . Let  $h(o_i) \in \mathbb{R}$  denote the exposure of offense  $o_i$  to reputational scrutiny, which allows the response to reputational pressure to vary across offenses. I define  $z_a$  and  $h(o_i)$  as deviations from their case-weighted means, so zero represents average exposure and the interaction terms capture differential responses relative to that benchmark. With this extension, a reviewer is characterized by

the pair  $(\alpha_a, z_a)$ , rather than just  $\alpha_a$ . The aggregate decomposition holds the exposure terms fixed and changes the case weights on the stable reviewer effects.

Define the extended office standard as

$$\tilde{s}_t^* = s_t^* + \eta_B B_t. \quad (12)$$

The first term is the incumbent's dynamic policy. The second allows reviewers to respond in common to supervision they expect from a current or anticipated principal. Because the common component  $g(t)$  depends on these terms only through their sum, its change over time cannot distinguish a change in the incumbent's policy from reviewers' common response to expected supervision. The office standard for reviewer  $a$  on case  $i$  becomes

$$\tilde{s}_{ait}^* = \tilde{s}_t^* + \lambda_B z_a B_t + \lambda_E h(o_i) E_t, \quad \lambda_B, \lambda_E \geq 0, \quad (13)$$

where  $z_a$  and  $h(o_i)$  are observed characteristics of the reviewer and the offense, respectively, and  $\lambda_B$  and  $\lambda_E$  are parameters. Substituting this office standard into the reviewer's best response gives

$$\tau_{ait}^* = (1 - w)\alpha_a + w\tilde{s}_{ait}^*. \quad (14)$$

Because I assume the heterogeneous exposure measures are centered on 0, the extension replaces  $g(t) = -\bar{f}ws_t^*$  with  $g(t) = -\bar{f}w\tilde{s}_t^*$  in Equation (11). It does not otherwise change the decomposition.

**Corollary 1 (Source-specific responses)** *In the extended model,*

$$\frac{\partial^2 \tau_{ait}^*}{\partial B_t \partial z_a} = w\lambda_B, \quad \frac{\partial^2 \tau_{ait}^*}{\partial E_t \partial h(o_i)} = w\lambda_E.$$

*These responses distinguish hierarchical from environmental pressure only when  $z_a$  modifies the response to hierarchy alone and  $h(o_i)$  modifies the response to reputational pressure alone.*

Approval in the extended model is

$$Y_{ait} = \mathbf{1}\{v_i > \tau_{ait}^*\} \mathbf{1}\{D_{it} = 0\}. \quad (15)$$

Other common forces are outside the state representation. If they moved during the period, the empirical office-wide component absorbs them and source attribution becomes weaker.

**Predictions.** Given the above, when no binding rule applies, the probability a reviewer approves felony charges in a given case is simply  $\Pr(Y_{ait} = 1 \mid x_i, a, t, D_{it} = 0) = 1 - G(\tau_{ait}^* \mid x_i)$ . The first thing to notice is that personnel and the common standard leave different traces. The decomposition attributes part of the change in office-wide approval to reviewer composition when case assignments shift toward reviewers with different (stable) approval propensities. A change in  $\tilde{s}_t^*$  instead enters the decisions of the reviewers who are active while preserving their baseline differences. Equation (11) separates the changing distribution of stable reviewer types from this common component. The dynamic model changes the possible sources of the common movement, but not this empirical distinction. If the common component changes, approval should also decline on average among reviewers observed on both sides of the transition.

**Implication 1 (Composition and common movement)** *Personnel change makes  $C_t^F$  dominate. A changing common standard makes  $C_t^s$  dominate and lowers approval among reviewers observed across the transition.*

Second, a categorical rule and the common standard have different consequences for reviewer decisions. A binding rule removes approval from the choice set for cases satisfying  $D_{it} = 1$ . In contrast, changes in the common standard shift the threshold for cases left to reviewer discretion. We can use categorical rules (i.e., the retail theft declination policy) as an empirical benchmark: the record should show a sharp response confined to the cases it covers. Thus, part of the empirical exercise entails establishing both the magnitude and the scope of policy-driven changes.

**Implication 2 (Categorical rule)** *A binding declination rule produces a sharp change among cases that satisfy its eligibility conditions. If eligibility is concentrated in an observed offense and the composition of cases reaching review remains stable, that offense shows a differential decline at the rule’s effective date. A changing common standard reaches discretionary decisions beyond that offense.*

Finally, notice that an abrupt change in approval does not by itself identify a categorical rule. An event that changes control of the office or beliefs about future control beyond prior expectations can also shift the office standard immediately. An anticipated transition can affect  $\tilde{s}_t^*$  earlier through the incumbent’s policy or reviewers’ common response to expected supervision. The incumbent may also change the standard before the election when office conduct affects her probability of remaining in office.

**Implication 3 (Onset)** *An electoral event that changes control or beliefs beyond prior expectations produces a level shift or a change in the subsequent path. The absence of such an effect does not reject electoral influence operating through the incumbent’s earlier adaptation or anticipation of a successor or environmental effects already under way.*

These implications organize the empirical analysis that follows. I evaluate whether the decline came from reviewer composition or common movement among active reviewers, whether the response to the retail-theft rule was concentrated among cases exposed to it, and whether political events changed the level or subsequent trajectory of approval. Together, these tests can identify the organizational mechanism and timing of change, but they cannot (by themselves) determine what caused any change in the common standard. The extended office standard,  $\tilde{s}_t^* = \sigma_{qt}(s_{t-1}^*, Z_t) + \eta_B B_t$ , combines the incumbent’s dynamic response to reputational and electoral conditions with reviewers’ common response to a current or anticipated principal. For example, the McDonald video changed the reputational environment and the incumbent’s electoral position at the same time. Foxx’s candidacy also changed beliefs about future control. Because these pressures all operated together, a change in  $\tilde{s}_t^*$  cannot separate their effects. Source-specific heterogeneity could distinguish them only if  $z_a$  modifies the response to hierarchical pressure alone and  $h(o_i)$  modifies the response to reputational pressure alone. However, there do not exist the kinds of *ex ante* covariates for reviewers or

offenses that support those exclusions, while other common pressures would weaken source attribution further. I therefore treat the source of the common movement as unresolved.

## 4 Data and Empirical Strategy

I evaluate the model’s predictions in the context of felony review in the Cook County State’s Attorney’s Office. In Cook County, police generally cannot file felony charges on their own.<sup>5</sup> After making a felony arrest, they call the State’s Attorney’s Felony Review Unit, where an on-call assistant state’s attorney usually evaluates the arrest within twenty-four hours. The reviewer approves felony charges, rejects them, or requests further investigation. Reviewers work fixed shifts, and incoming calls are assigned in rotation to the next reviewer on the active roster. Conditional on the roster and call time, the reviewer assigned to a case is therefore close to random with respect to its characteristics (Jordan 2024). The reviewer handles only this stage and then turns the case over to another attorney. Most reviewers are junior attorneys who rotate through the unit within a year or two.

The data here cover the period from January 2014 through December 2019. I chose this period because it spans a police scandal and a transition in leadership. In particular, in November 2015 the Chicago Police Department released a video of the police killing of Laquan McDonald a year earlier, following a court order to do so. The video contradicted official accounts of the killing and was central to subsequent murder charges against the officers who killed him. Shortly thereafter, in March 2016, Kim Foxx defeated the incumbent in the Democratic primary. She later won the general election and took office on December 1, 2016. The scandal and election were closely connected: the video that discredited the incumbent also reshaped the contest in Foxx’s favor. Two weeks later, her administration adopted a policy under which the office would not charge individuals for felony retail theft

---

<sup>5</sup>Illinois law permits police to file charges directly in narcotics cases, assuming no other felony is involved (Jordan 2025), and a recent experimental policy has expanded that authority to gun charges. Narcotics-only arrests account for roughly half of Chicago’s felony arrests during the period, so I exclude them from the population at risk of reaching review

unless the amount stolen was greater than \$1000, despite state law defining felony retail theft as stealing merchandise in excessive \$300 in value.

Specifically, I build a case-participant panel that follows individual cases from arrest through charging. I identify the universe of arrests using data from the Chicago Police Department, which includes the charges, governing statutes and offense classes, arrestee race, and arrest date. I link these records to the State’s Attorney’s databases that cover felony intake and each charge initiated against each arrestee. I link these data sources via the booking number and the State’s Attorney’s case-participant identifier, which I had to obtain via an Illinois Freedom of Information Act request. A separate crosswalk connects that identifier to the public case records. These links allow me to follow a Chicago arrest through felony review and charging. Roughly 95 percent of felony arrests match to a case record, with no meaningful trend from 2014 through 2018.<sup>6</sup>

Due to data availability, the data include only arrests made by the Chicago Police Department, the county’s dominant source of felony cases. Other municipal departments and the county sheriff also send cases to the State’s Attorney. (There are 249 agencies sending arrests to felony review during the period of study, 67.4% of which come from the Chicago PD.) As a result, for the analyses that only use review and charging record, I can use the full countywide caseload. Because the arrest data include arrestees’ race but not age or sex, arrest-side composition adjustments use race, offense, and charge class. Table 1 describes the countywide reviewed caseload before and after the transition. Figure 1 shows the full data funnel for all Chicago arrests I was able to link to prosecution data. This distinction between the countywide review sample and the linked Chicago sample carries through the analysis.

## 4.1 Measures

With these data, I build a set of variables on which the empirical strategy relies.

---

<sup>6</sup>The match rate declines modestly in late 2019, which I address in Appendix G. The unmatched share combines genuine non-prosecution with imperfect booking numbers, which I cannot fully separate.

Table 1: Sample description, by period. Countywide reviewed case-participants, 2014–2019.

|                            | Pre (2014–16) | Post (2017–19) |
|----------------------------|---------------|----------------|
| Reviewed case-participants | 59,813        | 51,031         |
| Black (%)                  | 61.0          | 66.3           |
| White (%)                  | 15.8          | 12.2           |
| Hispanic (%)               | 20.9          | 19.4           |
| Other (%)                  | 2.3           | 2.2            |
| Male (%)                   | 85.4          | 87.4           |
| Mean age at incident       | 33.2          | 32.8           |
| Felony-class charge (%)    | 99.8          | 99.9           |
| Felony-review approval (%) | 97.8          | 94.0           |

*Columns pool the 2014–16 and 2017–19 periods. The reviewed caseload is countywide; the arrest-linked funnel of Figure 1 covers Chicago Police Department arrests.*

**Felony-review approval.** This variable captures the central charging-stage outcome for each case-participant, measured at the level of the primary charge. I code a case as approved if its felony-review result is recorded as “Approved” or “Charge(S) Approved,” and as not approved if it is “Disregard,” “Rejected,” or “Continued Investigation.” I exclude observations where the felony review outcome is either missing or takes on an “other” value.

**Felony-eligible arrest and reaching review.** For the analysis that involves selection into felony review, I need the population of felony-eligible arrests before the felony review stage. I code an arrest as felony-eligible if any of its charges is recorded as a felony, and as *review-eligible* if at least one of its felony charges is not a narcotics offense, since narcotics-only arrests bypass review by rule.<sup>7</sup> I code an arrest as having reached felony review if it appears in the felony review database. The selection rate is the share of review-eligible arrests that reach review, conditioned on a successful arrest-to-case link.<sup>8</sup>

**Filed charge class.** I order Illinois offense classes from most to least severe, from the felony classes X, 1, 2, 3, and 4, with first-degree murder (Class M) ranked alongside X, then

<sup>7</sup>Narcotics statutes are 720 ILCS 570 (controlled substances), 550 (cannabis), 646 (methamphetamine), and 600 (drug paraphernalia).

<sup>8</sup>The conditioning removes the mechanical exclusions and the modest late-period drift in match quality from the measure.

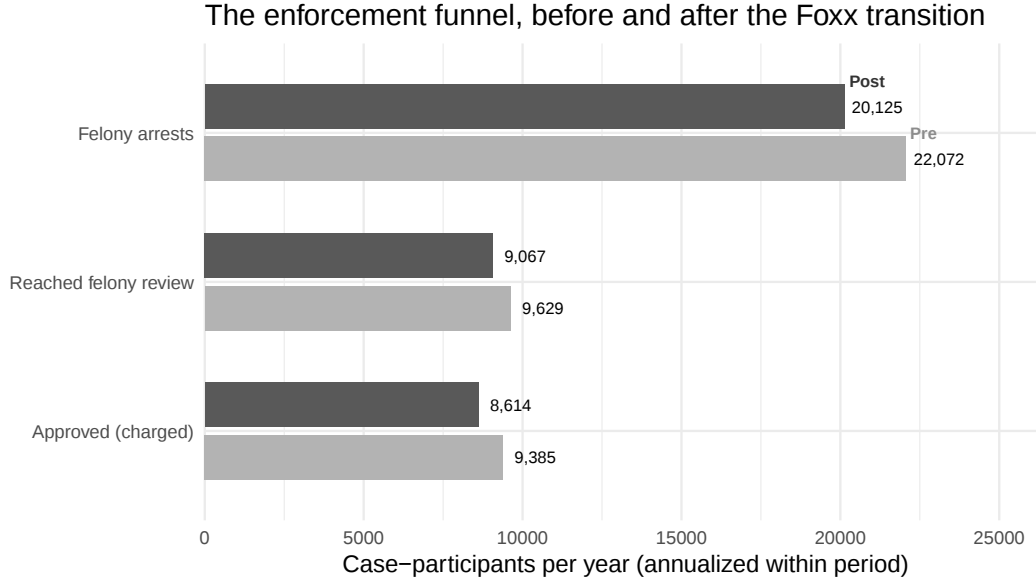


Figure 1: The enforcement funnel before (Pre, 2014–16) and after (Post, 2017–19) the Foxx transition: felony arrests, arrests reaching felony review, and approvals (charges filed). Counts are annualized and computed on the linked, Chicago-arrest-based funnel, so each stage is a subset of the one above it. Table 1 describes the full countywide reviewed caseload.

the misdemeanor classes A, B, and C. I use the class of the charge as filed at initiation to measure charge severity. Illinois codes all felonies into one of these classes, and 730 ILCS 5/5-4.5-10 maps each class directly into a sentencing range that constrains judges’ decisions.

**Reviewer type.** To distinguish changes in *who* reviews cases from changes in *how* reviewers decide, I estimate reviewer-specific fixed effects. These account for stable approval propensity, net of the common trends over time, the retail-theft policy, and observable case mix. The personnel contribution is the change in the case-weighted mean of these effects. Appendix C reports simpler leave-one-out and turnover measures as alternative diagnostics.

## 4.2 Empirical Strategy

The main empirical problem is that the cases prosecuted come from a population of arrests that the police, rather than the prosecutor, choose. Responding at least partially to its own broader political and institutional environment, police make the critical decision about which

arrests to make and present to prosecutors in the first instance (Richman 2003, Shjarback et al. 2017, Sklansky 2018). Thus, changes in aggregate charging rates may reflect changes in the standards the office uses to evaluate cases or changes in the mix of cases presented to the office (Knox, Lowe and Mummolo 2020, Ouss and Rappaport 2020). Using the universe of arrests lets me observe the population before felony review, so I can describe how the case pool changed and bound whether that change can account for the movement in approval.<sup>9</sup> I present the analysis in three steps, corresponding to Implications 1, 2, and 3. The first step also asks whether changes in the reviewed caseload can account for the estimated common movement.

The empirical analysis rests on a simple baseline specification. Following the model, I capture the office-wide movement with a smooth term in a linear probability model of the review decision,

$$\mathbb{E}[\text{approve}_{it} \mid x_{it}, a(i), t] = g(t) + \gamma \text{policy}_{it} + x'_{it}\beta + \underbrace{\tilde{\alpha}_{a(i)}}_{\text{type}} + \underbrace{\nu_{a(i)} t}_{\text{within-reviewer deviation}}, \quad (16)$$

where  $g(\cdot)$  is a flexible office-wide time component,  $\text{policy}_{it}$  indicates exposure to the retail-theft policy,  $x_{it}$  contains case covariates, and  $\tilde{\alpha}_a$  is a reviewer fixed effect for reviewer  $a(i)$ . Equation (10) gives the specification its threshold interpretation, and each term has a counterpart in the model. The time component is the approval-scale image of the common standard,  $g(t) = -\bar{f}w\tilde{s}_t^*$ . The reviewer effect is the approval-scale reviewer type  $\tilde{\alpha}_a = -\bar{f}(1-w)\alpha_a$ , which the model holds fixed over time. The policy term stands in for the categorical rule  $D_{it}$ . The reviewer-specific slope  $\nu_a$  allows a reviewer to depart from the common movement.<sup>10</sup>

I estimate a series of variants of Equation (16). First, I use Equation (11) to separate

---

<sup>9</sup>I do not identify why police behavior changed or whether it responded to the prosecutor, a subject that deserves more theoretical and empirical attention (cf. Clark 2026a, Ouss and Rappaport 2020, Richman 2003).

<sup>10</sup>Appendix B reports the institutional basis for the assignment claim and balance tests. Assignment is determined by the active roster of felony review attorneys and the timing of incoming calls, and observed case attributes explain little of which reviewer receives a case.

changes in review rates into reviewer composition and movement in the common standard. I then examine approval changes among reviewers observed before and after the leadership transition. The second part uses the December 15, 2016 retail-theft policy as a benchmark for the prediction that a binding rule produces a sharp change confined to the cases it covers (Implication 2). Because the policy took effect two weeks after Foxx’s inauguration, a retail-specific response that holds common calendar movement apart distinguishes the formal rule from an office-wide change at the transition. I then turn to Implication 3, which concerns the shape of  $g(t)$  at key political moments. I estimate the model to assess whether we see shifts in either the rate of approval or the trend in the rate of approval associated with the critical events discussed above. Importantly, failing to detect a break does not reject electoral influence operating through the incumbent’s adaptation before the election, through anticipation of a successor, or through gradual movement too small to separate from a change already under way. Thus, I focus on establishing what magnitudes it can rule out.

## 5 Where Did Change Occur?

### 5.1 Implication 1: The decline appears as common movement

The first empirical question I address is whether the office changed because different reviewers handled its cases or because active reviewers changed how they decided them. It is possible to observe the same aggregate decline because of each mechanism, but for different reasons. Implication 1 identifies how we can distinguish them in the data. Personnel change shifts cases toward reviewers with different stable approval propensities, making  $C_t^F$  dominate Equation (11). The common standard mechanism changes approval within the active set of reviewers, making  $C_t^s$  dominate.

For the decomposition, I set  $\nu_a = 0$  and estimate the following specification:

$$\mathbb{E}[\text{approve}_{it} \mid a(i), t, x_{it}] = g(t) + \gamma \text{policy}_{it} + x'_{it}\beta + \tilde{\alpha}_{a(i)}, \quad (17)$$

where  $g(t) = \beta_0 + \sum_{k=1}^{12} \theta_k B_k(t)$  and  $B_k(t)$  are the twelve basis functions of a natural spline in review month. The policy indicator is

$$\text{policy}_{it} = \mathbf{1}\{o_i = \text{Retail Theft}\} \mathbf{1}\{t \geq \text{January 2017}\}.$$

This offense-level indicator proxies for exposure to the categorical rule  $D_{it}$ . The vector  $x_{it}$  contains linear controls for age and charge severity and fixed effects for race, sex, and broad offense type. The reviewer effect is  $\tilde{\alpha}_{a(i)}$ . This specification is Equation (16) with  $\nu_a = 0$  and  $g(t)$  represented by the natural spline.

The reviewer fixed effects capture that stable approval propensity within reviewers. The personnel contribution,  $C_t^F$ , is the change between 2014 and 2019 in the case-weighted mean of those effects. This allows for a change in the share of cases handled by each reviewer while holding each reviewer’s baseline approval propensity fixed. The common contribution,  $C_t^s$ , is the change in  $g(t)$  over the same period. It captures movement shared across the office. The two components sum to the adjusted change in approval.

As we saw in Table 1, felony-review approval fell by about four percentage points over the window, from an average of 97.8 percent in 2014–16 to 94.0 percent in 2017–19. Figure 2 plots the common time component estimated from Equation (17). I estimate the model using 93,378 decisions by 321 reviewers who handled at least 25 cases. Each gray point is the monthly approval rate after removing the estimated contributions of reviewer identity, the retail-theft policy, and observed case mix. The solid line is the fitted spline under the same normalization.<sup>11</sup>

The fitted line falls 4.27 percentage points between January 2014 and December 2019, dropping from 98.21% to 93.95%. Table 2 reports compares case-weighted annual averages; by that measure, the common component is 3.29 percentage points lower in 2019 than in

<sup>11</sup>Specifically, each point is the fitted spline for that month plus the mean regression residual among cases reviewed during the month. I shift both series by a single constant so that their case-weighted means equal the sample approval rate. The absolute vertical level is therefore a normalization; only differences in  $g(t)$  over time are identified.

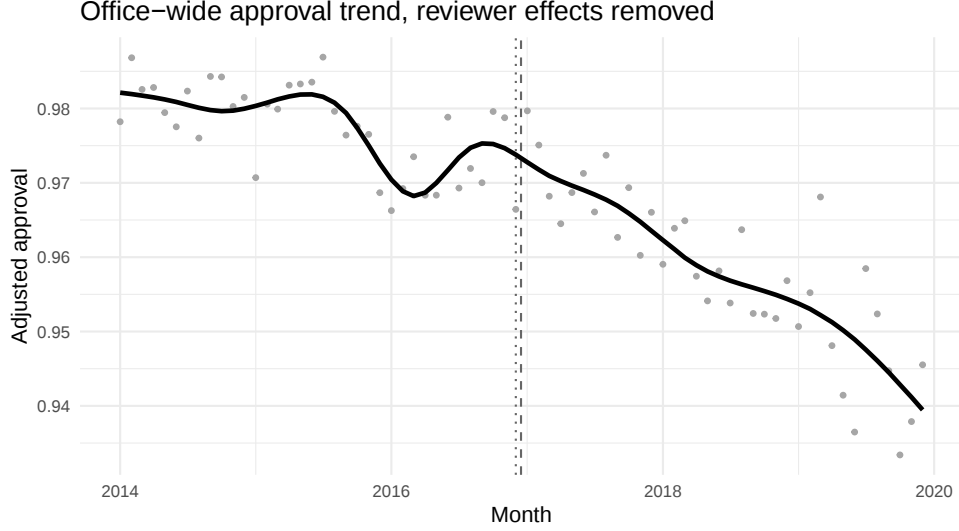


Figure 2: The office-wide calendar component of felony-review approval. The gray point for month  $t$  is  $\hat{\kappa} + \hat{g}(t) + \hat{\epsilon}_t$ , where  $\hat{\epsilon}_t$  is the mean residual from the case-level model described in the text. The solid line is  $\hat{\kappa} + \hat{g}(t)$ . The constant  $\hat{\kappa}$  returns the series to the approval-rate scale; its level is a normalization. The dotted line marks Foxx’s inauguration and the dashed line the retail-theft rule.

2014. Of course, the temporal component,  $g(t)$ , may reflect common movement due to the incumbent’s policy or a common response to expected supervision, so these estimates do not capture the source of change.

Table 2 breaks down the two components. Because cases are essentially randomly assigned to reviewers conditional on date and time, the estimated reviewer effects can be interpreted as differences in reviewer type rather than differences in the cases assigned to them (see also Jordan 2024). Holding case mix and the retail-theft policy constant, approval declined by 4.05 percentage points between 2014 and 2019 (95 percent confidence interval  $[-4.58, -3.51]$ ). Reviewer composition accounts for a 0.76-point decline, or 18.7% of the total. However, the confidence interval for that estimate ranges from a 2.08-point decline to a 0.57-point increase. The common component accounts for the remaining 3.29 points, or 81.3 percent, and is more precisely estimated (95 percent confidence interval  $[-4.69, -1.89]$ ). The point estimate therefore assigns about four times as much of the decline to common movement as to reviewer composition.

Table 2: Additive decomposition of the change in felony-review approval.

| Quantity                                       | Estimate | 95% CI         |
|------------------------------------------------|----------|----------------|
| <i>Additive OLS decomposition, 2014–2019</i>   |          |                |
| Adjusted approval change (percentage points)   | −4.05    | [−4.58, −3.51] |
| Reviewer composition $C^F$ (percentage points) | −0.76    | [−2.08, 0.57]  |
| Common movement $C^s$ (percentage points)      | −3.29    | [−4.69, −1.89] |
| Composition share (percent)                    | 18.7     | [−13.9, 51.3]  |
| Common-movement share (percent)                | 81.3     | [48.7, 113.9]  |

*The table implements Equation (11) in a linear probability model using reviewers with at least 25 cases. Reviewer types are fixed effects, and a twelve-degree-of-freedom natural spline measures the common time component. The 2014 and 2019 components use the actual case weights in each year, with inactive reviewers assigned zero weight. Confidence intervals cluster by review month. The reviewer–month graph forms a single connected component. Appendix C reports sensitivity to the caseload threshold and flexibility of the common trend.*

Turning from the office-wide, common change in approval, I evaluate change in approval among a stable group of reviewers. Ideally, I would estimate the distribution of the reviewer-specific slopes  $\nu_a$ . Unfortunately, reviewer tenures are too short and uneven to estimate that distribution precisely. I instead restrict the sample to reviewers who made at least ten decisions on each side of the December 2016 inauguration and estimate annual changes relative to 2016. This holds the membership of the group fixed by construction. Among the 82 reviewers who meet that requirement, adjusted approval in 2019 is 3.61 percentage points ([−6.87, −0.35]) below its 2016 level (see Figure 3). Requiring between five and twenty decisions on each side produces declines between 3.45 and 3.90 points. However, it is important to underscore that this does not imply that every reviewer moved in parallel. Rather, it simply shows that the decline was not confined to replacing reviewers with less permissive ones.

Taken together, the point estimates attribute most of the decline to common movement rather than reviewer composition. The spanning-reviewer estimates also show that the decline was not confined to turnover. These results do not identify the source of the common change; they are consistent with adaptation by the incumbent, anticipation of her successor, or a change in the office’s reputational environment.

Of course, we might be concerned that the office changed how it screened cases. First, the

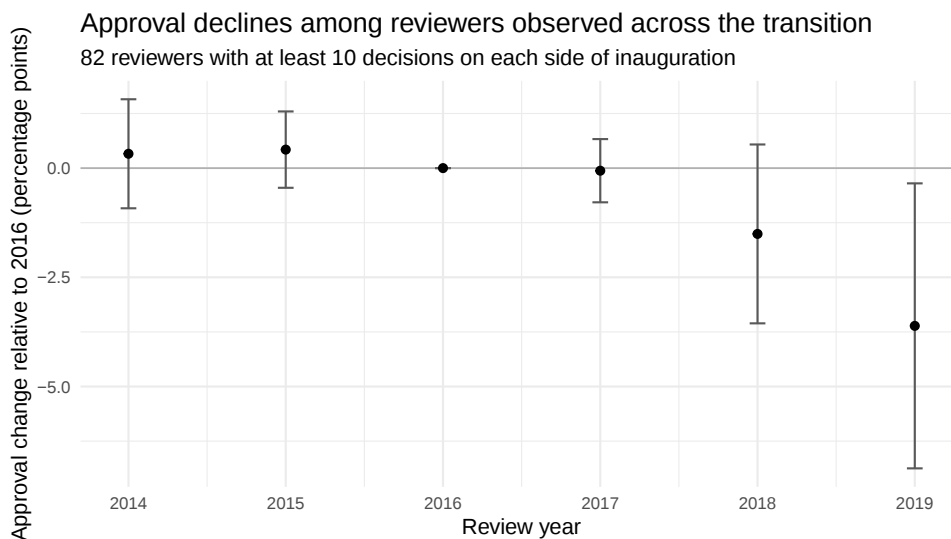


Figure 3: Annual approval changes among reviewers with at least ten decisions on each side of the December 2016 inauguration. Estimates are relative to 2016 and come from a linear probability model with reviewer fixed effects, case controls, and the retail-theft policy indicator. Bars report 95 percent confidence intervals clustered by reviewer.

office might have adopted a new pattern of screening across cases rather than changed the level of the standard it applied. However, there is no evidence for a new screening pattern. As I show in Appendix H, the share of the reviewed caseload made up of offenses in which investigation reliably produces evidence is nearly flat, and the decline is similar for those offenses and for offenses that depend more heavily on testimony. Case characteristics also explain less, rather than more, of the approval decision over time. The change is therefore broad rather than concentrated along an observable dimension of case strength.

Second, the cases reaching felony review changed over the period. Chicago felony arrests fell from about twenty-five thousand in 2014 to about eighteen thousand in 2016–17, then rebounded to about twenty-two thousand by 2019. Among review-eligible arrests, the share reaching review fell from about 0.97 in 2014 to about 0.90 in 2019. Recalculating each year’s approval rate using the 2014 proportions of cases in each race-by-felony-class category produces nearly the same annual trend. Moreover, while the approval rate fell over the period, tighter selection into review should, if anything, leave the office with stronger cases and push approval upward. Under the monotonicity assumption described in Appendix F,

trimming bounds place the within-pool change below zero from 2016 forward. The covariate-tightened bound for 2018 is approximately  $[-0.036, -0.027]$ , and its Imbens–Manski (Imbens and Manski 2004) confidence interval also remains below zero. Collectively, these robustness checks suggest that the decline in approval is not due to changes in the cases presented to the office.

## 5.2 Implication 2: A categorical rule produces an offense-specific break

Turning to Implication 2, the question is whether we can distinguish the effects of a categorical rule affecting felony review for some cases from an office-wide change in review thresholds. On December 1, 2016, Foxx took office, and two weeks later announced a formal policy on felony retail theft declination. Illinois law at the time classified retail theft above \$300 as a felony, while ordinary theft carried a \$500 threshold. Effective December 15, 2016, the office presumptively declined to charge retail theft as a felony unless the value reached \$1,000 or the accused had at least ten prior felony convictions. However, because the policy was announced just 2 weeks after Foxx took office, any change in felony review approval rates around that policy could therefore combine a targeted policy response with office-wide movement. Thus, I rely on the logic from Implication 2, which highlights that the rule should produce an immediate response, but only among the cases it covers.

Unfortunately, the data do not consistently report the value of the stolen property or the defendant’s criminal history, so I cannot identify every case for which the rule bound (i.e., retail theft between \$300 and \$1000). I instead estimate an offense-level intention-to-treat response by comparing all retail-theft cases with ordinary-theft cases. This comparison mixes retail-theft cases covered by the policy with cases that remained eligible for felony prosecution. Thus, the estimates I report here reflect the share exposed and compliance with the rule. However, because the policy also changed which cases reached review, it need not be an attenuated version of the effect among covered cases. (It is possible, for

example, that policy simply stopped referring cases for review when they were covered by the declination policy.) In this sense, the estimates here constitute a reduced-form response for retail theft as an offense category rather than the structural effect of  $D_{it} = 1$ .

I estimate a simple specification that modifies the baseline model to an event-study. This specification allows the retail-theft differential to vary across periods. It shows whether the offenses were already diverging before the rule and whether the differential appears when the rule takes effect. Formally, I replace  $g(t)$  with calendar-month fixed effects, restrict the sample to retail and ordinary theft, and estimate

$$\mathbb{E}[\text{approve}_{it} \mid a(i), t, x_{it}] = \mu_t + \rho R_i + \sum_{q \neq -1} \gamma_q R_i \mathbf{1}\{Q_{it} = q\} + x'_{it}\beta + \tilde{\alpha}_{a(i)}. \quad (18)$$

where  $\mu_t$  is a calendar-month fixed effect,  $R_i$  indicates retail theft, and  $x_{it}$  contains linear controls for age and charge severity and fixed effects for race and sex. The reviewer effect is  $\tilde{\alpha}_{a(i)}$ . September through November 2016 form the omitted period,  $q = -1$ . December 2016 is a separate transition period,  $q = 0$ , because the rule took effect halfway through the month. The remaining periods pool three months each, and the estimation window covers two years on either side of the policy. Standard errors are clustered by calendar month.

Figure 4 reports the estimates covering two years on either side of the policy. It groups months into three-month periods, except that December 2016 remains a separate transition period because the rule took effect halfway through the month. September through November 2016 form the omitted reference period. Before the adoption of the declination policy, retail and ordinary theft follow similar paths: the seven prepolicy coefficients are jointly indistinguishable from zero ( $p = .56$ ). In December, retail-theft approval falls by 6.11 ( $[-9.23, -2.98]$ ) percentage points relative to ordinary theft and the prepolicy difference. The response is therefore confined to the named offense and begins in the month the rule takes effect.

Pooling the full post-policy period produces the same conclusion: retail-theft approval

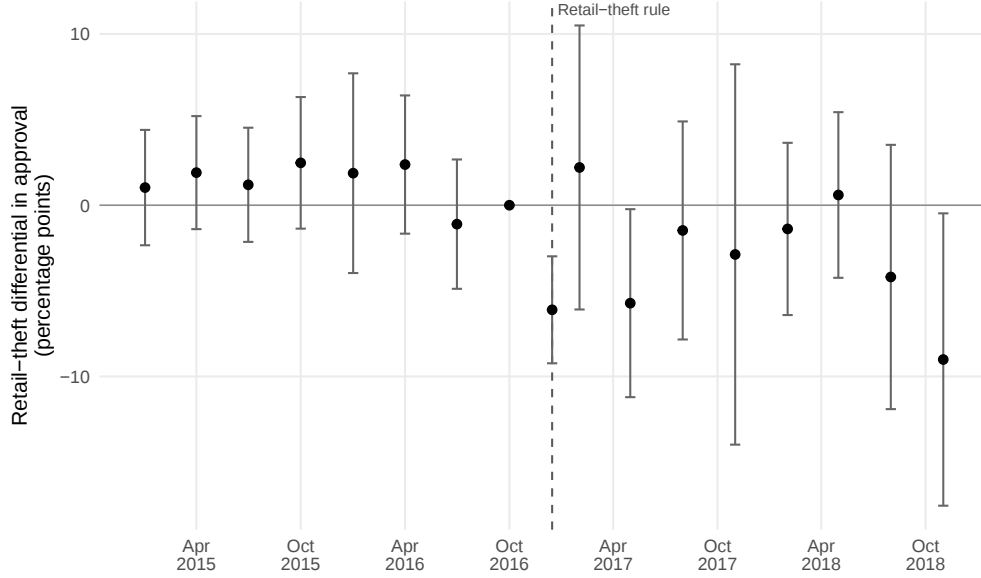


Figure 4: Retail-theft approval around the December 15, 2016 declination policy, relative to ordinary theft. Points are differences from the September–November 2016 retail-theft differential. December 2016 is its own transition period; the remaining estimates pool three months. Bars report 95 percent confidence intervals clustered by calendar month.

fell by 3.56 ( $[-5.77, -1.34]$ ). In Appendix D, I report the full results of this specification and consider others, such as employing December 15 as the treatment date and using alternative comparison groups. The December drop does not persist in every subsequent quarter, but the set of retail-theft cases reaching review also changes sharply. Retail theft falls from about 15% of the reviewed docket to about 5% within weeks and remains there through 2019 (Appendix F). The cases we observe after the policy was adopted therefore mix new ones with pre-policy cases.

The key finding here is the contrast between the change in approval of cases (potentially) covered by the declination policy and the temporal trend in approval. Whereas the policy leads to an immediate change for the offense it names, the office-wide path remains smooth after holding that response apart. In this way, the results here distinguish an immediate, offense-specific change (of this magnitude) from an office-wide change in the common standard.

### 5.3 Implication 3: The decline predates the political transition

Finally, I turn to Implication 3, which concerns when the office-wide change in approval began. As the model illustrates, a variety of potential events might cause either an immediate shift in approval or a change in its subsequent path, but any given event need not produce such a change. For example, incumbents may adapt before voters act, and reviewers may respond once a successor becomes foreseeable. In Cook County, Foxx entered the race by the late summer of 2015, the McDonald video was released that November, she won the binding Democratic primary in March 2016, and she took office in December. Here, I estimate a series of specifications that assess whether any of these events is associated with a detectable change in the change that may already have been underway.

The specification I estimate involves setting  $\nu_a = 0$  for a candidate month  $c$  and replacing the smooth component in Equation (16) with a segmented linear path:

$$\begin{aligned} \mathbb{E}[\text{approve}_{it} \mid a(i), t, x_{it}] = & \beta_0 + \beta_1 T_{tc} + \delta_c H_{tc} + \kappa_c T_{tc}^+ \\ & + \gamma \text{policy}_{it} + x'_{it} \beta + \tilde{\alpha}_{a(i)}, \end{aligned} \tag{19}$$

where  $T_{tc}$  is the number of months from candidate month  $c$ ,  $H_{tc} = \mathbf{1}\{T_{tc} \geq 0\}$ , and  $T_{tc}^+ = \max\{T_{tc}, 0\}$ . The coefficient  $\delta_c$  captures the immediate level shift at  $c$ , while  $\kappa_c$  estimates the change in the monthly slope after  $c$ . The change relative to the pre-event path after one year is  $\delta_c + 12\kappa_c$ . The policy indicator is the retail-theft-by-January 2017 interaction defined below Equation (17);  $\gamma$  holds that response apart from the office-wide path. The vector  $x_{it}$  contains linear controls for age and charge severity and fixed effects for race, sex, and broad offense type. I also include reviewer fixed effects. The joint test of a break is  $H_0 : \delta_c = \kappa_c = 0$ .

I use this segmented path in two ways. First, I allow the data to select the break date. I remove the estimated contributions of reviewer identity, the retail-theft policy, and observed case mix from approval and calculate the adjusted mean for each month. I fit the segmented

path at every interior month, weighting each monthly observation by its number of cases, and select the month that minimizes the residual sum of squares. I estimate standard errors by bootstrapping the fitted model’s monthly residuals in six-month blocks and repeating the break-date search for 1,000 resampled series. Second, I estimate Equation (19) for the McDonald video, the Democratic primary, and Foxx’s inauguration over the full 2014–2019 period. I code the video and primary variables using the first full month after each event, December 2015 and April 2016. I code the inauguration using December 2016. I aggregate model scores by review month and use a Newey–West covariance estimator with a six-month lag.<sup>12</sup>

The estimated unrestricted break occurs in September 2015; however the confidence interval ranges from February 2015 to September 2018 (Figure 5). Thus, the point estimate for the onset of change in approval rates precedes the video, primary, and inauguration, but the interval contains all three. The data-driven break therefore suggests that the decline began before the formal transition without localizing its onset to a single political event.

However, the candidate-date tests yield more precise estimates. In the first full month after the McDonald video’s release, adjusted approval falls by 0.59 percentage points ( $[-0.94, -0.24]$ ). Combining the immediate shift with the change in slope yields a 1.03-point decline after one year ( $[-1.70, -0.36]$ ), and the two terms are jointly different from zero ( $p = .001$ ).<sup>13</sup> Because the video changed both the office’s reputational environment and the incumbent’s electoral position, these estimates cannot separate those two pressures. At the same time, neither the primary nor the inauguration is associated with a comparable break. The primary is associated with a +0.16 percentage point ( $[-0.40, 0.72]$ ) shift, and the estimated change after one year is  $-0.18$  points ( $[-1.11, 0.74]$ ). (The joint  $p$ -value is .55). The immediate change associated with the inauguration is +0.14 points ( $[-0.50, 0.78]$ ), and the change after one

---

<sup>12</sup>A null result at a candidate date can reflect either the absence of a break or limited precision. I therefore supplement the joint tests with a full grid of placebo dates, shorter estimation windows, equivalence tests against a one-percentage-point bound, and minimum detectable effects at 80% power.

<sup>13</sup>The video ranks third among the 48 admissible dates in the placebo grid, and the result persists in twelve-, eighteen-, and twenty-four-month windows.

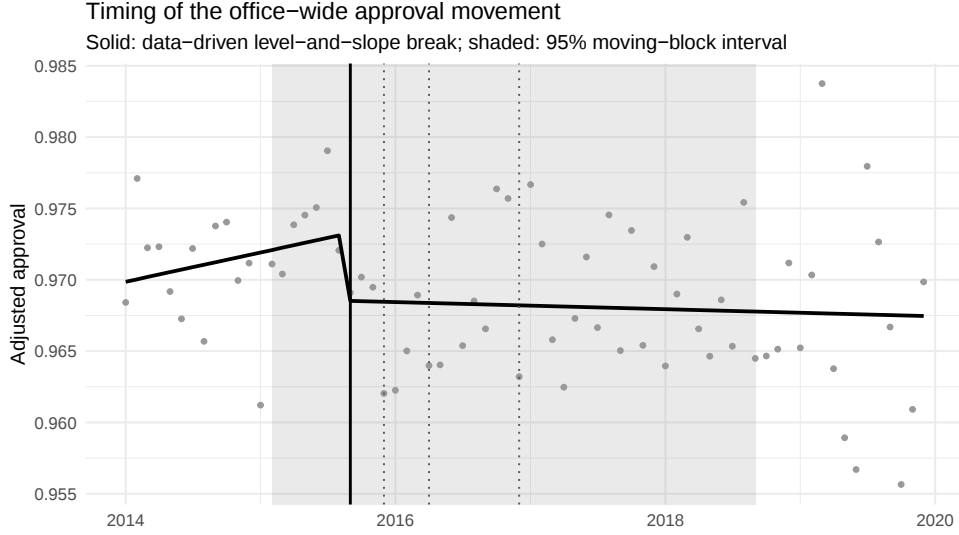


Figure 5: Timing of the office-wide approval movement. The solid line shows the least-squares level-and-slope break in the adjusted monthly series. The shaded band is its 95 percent moving-block bootstrap interval, and the dotted lines show the McDonald video, Democratic primary, and inauguration.

Table 3: Candidate-date changes in felony-review approval. Entries are percentage-point estimates with 95 percent confidence intervals. Each row is a case-level linear probability model with reviewer fixed effects, case controls, the retail-theft policy indicator, and a full-period segmented linear trend. The annual slope change is twelve times the monthly post-date slope coefficient. The final column jointly tests the level and slope changes. Newey–West standard errors aggregate scores by month and allow six months of serial dependence.

| Candidate date     | Immediate shift      | Annual slope change | Change after one year | Joint $p$ |
|--------------------|----------------------|---------------------|-----------------------|-----------|
| McDonald video     | −0.59 [−0.94, −0.24] | −0.44 [−1.04, 0.16] | −1.03 [−1.70, −0.36]  | 0.001     |
| Democratic primary | 0.16 [−0.40, 0.72]   | −0.35 [−1.08, 0.39] | −0.18 [−1.11, 0.74]   | 0.553     |
| Inauguration       | 0.14 [−0.50, 0.78]   | −0.68 [−1.38, 0.02] | −0.54 [−1.41, 0.33]   | 0.162     |

year is  $-0.54$  points ( $[-1.41, 0.33]$ ); the joint  $p$ -value is .16 (Table 3). For both the video and the primary, equivalence tests rule out immediate shifts of at least one percentage point. The estimates therefore rule out large immediate responses at the electoral dates more decisively than moderate gradual influence after Foxx took office.<sup>14</sup>

Together, these findings suggest that the change in felony review had begin before both

<sup>14</sup>The full-period design has 80 percent power to detect an immediate shift of about 0.91 points at inauguration, but a gradual one-year change would have to reach about 1.24 points to be detected with the same power. Eighteen- and twenty-four-month windows also show no joint break at either date. Twelve-month windows fluctuate because they contain adjacent events; Appendix E reports the full sensitivity analysis.

the primary and inauguration. Nor is there evidence that the trend changed with either event. That does not rule out electoral influence, of course, as the incumbent could have adapted in response to electoral pressure or because Foxx’s candidacy made a different future principal foreseeable. In short, the McDonald video had implications for both the State’s Attorney’s office’s reputational environment and the electoral campaign that would determine its future leadership. Because the data contain no predetermined exposure measure that separates these pressures, what we have is an estimate of change that predates the formal transition, not a causal estimate of which political force was driving it in the first place.

## 6 Discussion and Conclusion

Taken together, the results here emphasize and clarify often overlooked theoretical aspects of standard political science accounts of the politics of bureaucratic control. In particular, they underscore the notion that a bureaucracy can change before political authority formally changes hands. Felony-review approval in Cook County declined mainly because the reviewers already working in the office became less likely to approve charges. The data here indicate that changes in personnel and the cases they handle play a much smaller role in driving the decline in felony approval. Of course, Foxx’s retail-theft rule produced the sharp response expected for the cases it directly targeted, but the broader decrease in felony approval was already under way when she won the primary and took office. Foxx therefore inherited an office that had begun to change through the decisions of incumbent reviewers.

This finding has deep implications for theories of political control of bureaucratic decision-making. First, consistent with myriad theories of agency and oversight, personnel authority can operate without personnel replacement. A line official who expects closer supervision or reassignment has reason to adapt before the principal removes anyone. However, in contrast to most empirical and theoretical analyses of principal-agent control of a bureaucracy, these findings also highlight the ways in which the same logic applies when career officials anticipate

the principal they expect to serve. There is, for example, no clear evidence that hierarchical pressure led to broad, office-wide change in Cook County. At the same time, the evidence suggests that the absence of turnover is not necessarily evidence that the office remained unchanged. Instead, these findings highlight the idea that political control can operate through the discretion that incumbent officials retain.

The distinction between personnel authority and personnel replacement also bears on debates over whether removal authority secures democratic control of administration.<sup>15</sup> Observed replacement is an incomplete measure of control because the prospect of personnel action may change conduct among those who remain. This implication extends beyond elected prosecutors to bureaucracies in which political leaders inherit officials who exercise substantial discretion. Indeed, the central contribution here is to show that formal succession can be a lagging indicator of bureaucratic change. An incumbent facing an electoral threat may adapt before voters act, just as career officials may respond to an anticipated successor. And, when external events that shape the reputational environment surrounding the bureaucracy occur, both pressures can induce bureaucratic change, independent of and before an actual transfer of power to a new political principal. In this sense, office-taking is a part of a broader political process of change, rather than the source of it.

This creates an identification problem. Analyses that treat political transitions as empirical indicators of change may describe the difference associated with a new political period, but it does not necessarily isolate the incoming leader's contribution. The pre-inauguration period may already contain adaptation by the incumbent or anticipation among career officials. This is a general lesson for any study of bureaucratoc control but it has particular

---

<sup>15</sup>This bears on a premise the Supreme Court sharpened in *Trump v. Slaughter* (2026), which overruled *Humphrey's Executor v. United States*, 295 U.S. 602 (1935), and held that the President may remove the heads of independent agencies at will, resting in part on the view that control over personnel secures democratic control over administration. The companion case *Trump v. Cook* (2026), decided the same day, preserved the Federal Reserve's independence, which the Court treated as distinguishable. A county prosecutor's office is not a federal independent agency, and one case cannot resolve the constitutional question. But the finding here is a reminder that formal authority over personnel is not the same as observable replacement: removal power can shape conduct through the threat it poses to those who stay, and the within-reviewer movement I document is consistent with such an incentive effect rather than with replacement as the mechanism of change.

bite in the context of the politics of criminal justice, an area bureaucracy of exceptionally high stakes. In Cook County, the release of the McDonald video changed the CCSAO’s reputational environment while simultaneously weakening the incumbent’s electoral position. That means that we cannot empirically distinguish a reputational response from reviewers’ adapting their behavior in anticipation of Foxx’s election. That does not mean that Foxx was irrelevant, of course. However, the evidence does show that neither the primary nor her inauguration identifies the onset of the office-wide change or her contribution to that change. This observaiton, in turn, has direct implications for various empirical studies, including the growing literature focused on evaluations of criminal justice and prosecutorial reform. Studies centered on an the election of a new prosecutor various estimate how imprisonment or charging changed across political regimes (Agan et al. 2025, Gottlieb et al. 2025). That is a consequential estimand, but it differs from the effect attributable to the incoming prosecutor after accounting for adaptation that preceded her. Prosecutorial outcomes also arise from a selected enforcement pipeline. Police determine which cases reach the prosecutor, and there are reasons to suspect that changes in prosecutorial policies may shape the set of cases that police present to a prosecutor in the first instance (Clark 2026b, Goldrosen 2026). Thus, combining data on policing and arrests with prosecutorial charging decisions can help separate changes in the cases presented to a prosecutor from changes in how prosecutors use their discretion (Knox, Lowe and Mummolo 2020, Ouss and Rappaport 2020). This is an important avenue for continuing research on the politics of justice and law enforcement, beyond bureaucratic control.

## References

- Agan, Amanda, Jennifer L. Doleac and Anna Harvey. 2023. “Misdemeanor Prosecution.” *Quarterly Journal of Economics* 138(3).
- Agan, Amanda Y., Jennifer L. Doleac, Anna Harvey, Anna Kyriazis and Lauren Schechter.

2025. Prosecutorial Reform and Local Crime Rates. NBER Working Paper 34364 National Bureau of Economic Research.
- Birkland, Thomas A. 1998. "Focusing Events, Mobilization, and Agenda Setting." *Journal of Public Policy* 18(1):53–74.
- Bolton, Alexander, John M. de Figueiredo and David E. Lewis. 2021. "Elections, Ideology, and Turnover in the U.S. Federal Government." *Journal of Public Administration Research and Theory* 31(2).
- Bowman, Rachel A., Belén Lowrey-Kinberg and Jon B. Gould. 2026. "Progressive Prosecution in Context: Examining the Impact of Prosecutorial Administration Change on Case-Processing Patterns and Racial Disparities." *Criminology & Public Policy* . Published online.
- Brehm, John and Scott Gates. 1997. *Working, Shirking, and Sabotage: Bureaucratic Response to a Democratic Public*. University of Michigan Press.
- Bustos, Edgar O. 2021. "Organizational reputation in the public administration: A systematic literature review." *Public Administration Review* 81(4):731–751.
- Busuioc, E Madalina and Martin Lodge. 2016. "The reputational basis of public accountability." *Governance* 29(2):247–263.
- Carpenter, Daniel P. 2010. *Reputation and Power: Organizational Image and Pharmaceutical Regulation at the FDA*. Princeton University Press.
- Clark, Tom S. 2026a. "Coordinated Justice: The Politics of Prosecution and Policing." Working paper, Stanford University.
- Clark, Tom S. 2026b. "A Model of Criminal Justice Discretion." Working paper, Stanford University.

- Clark, Tom S. forthcoming. “Stacking the Charges: Prosecutorial Discretion in Criminal Charging Decisions.” *American Journal of Political Science* .
- Doherty, Kathleen M., David E. Lewis and Scott Limbocker. 2019. “Controlling Agency Choke Points: Presidents and Regulatory Personnel Turnover.” *Journal of Public Administration Research and Theory* 29(3).
- Dunlea, RR and Don Stemen. 2026. “The impact of prosecutors’ office caseloads on case processing outcomes.” *Criminology & Public Policy* 25(2):297–322.
- Gailmard, Sean and John W. Patty. 2012. “Formal Models of Bureaucracy.” *Annual Review of Political Science* 15:353–377.
- Gilad, Sharon, Moshe Maor and Pazit Ben-Nun Bloom. 2015. “Organizational Reputation, the Content of Public Allegations, and Regulatory Communication.” *Journal of Public Administration Research and Theory* 25(2):451–478.
- Goldrosen, Nicholas. 2026. “The Interorganizational Effects of Reform Prosecution Declination Policies.” *Justice Quarterly* . Published online.
- Gordon, Sanford C. and Gregory A. Huber. 2002. “Citizen Oversight and the Electoral Incentives of Criminal Prosecutors.” *American Journal of Political Science* 46(2).
- Gottlieb, Aaron, Matt Epperson, Laura Giugno and Hye-Min Jung. 2025. “What happens to imprisonment rates when a progressive prosecutor is elected: Quasi-experimental evidence from Cook County.” *Punishment & Society* 27(2):315–337.
- Huber, John D. and Charles R. Shipan. 2002. *Deliberate Discretion? The Institutional Foundations of Bureaucratic Autonomy*. Cambridge University Press.
- Imbens, Guido W. and Charles F. Manski. 2004. “Confidence Intervals for Partially Identified Parameters.” *Econometrica* 72(6):1845–1857.

- Jordan, Andrew. 2024. “Racial Patterns in Approval of Felony Charges.” Working paper, Washington University in St. Louis.
- Jordan, Andrew. 2025. “Relaxing Charging Restrictions in Narcotics Cases.” Working paper, Washington University in St. Louis.
- Kingdon, John W. 1984. *Agendas, Alternatives, and Public Policies*. Little, Brown.
- Kline, Patrick, Raffaele Saggio and Mikkel Sølvsten. 2020. “Leave-Out Estimation of Variance Components.” *Econometrica* 88(5):1859–1898.
- Knox, Dean, Will Lowe and Jonathan Mummolo. 2020. “Administrative Records Mask Racially Biased Policing.” *American Political Science Review* 114(3).
- Lee, David S. 2009. “Training, Wages, and Sample Selection: Estimating Sharp Bounds on Treatment Effects.” *Review of Economic Studies* 76(3):1071–1102.
- Lerman, Amy E. and Vesla M. Weaver. 2014. *Arresting Citizenship: The Democratic Consequences of American Crime Control*. University of Chicago Press.
- Lewis, David E. 2008. *The Politics of Presidential Appointments: Political Control and Bureaucratic Performance*. Princeton University Press.
- Manski, Charles F. 1990. “Nonparametric Bounds on Treatment Effects.” *American Economic Review* 80(2):319–323.
- McCubbins, Mathew D, Roger G Noll and Barry R Weingast. 1987. “Administrative procedures as instruments of political control.” *The Journal of Law, Economics, and Organization* 3(2):243–277.
- McCubbins, Mathew D. and Thomas Schwartz. 1984. “Congressional Oversight Overlooked: Police Patrols versus Fire Alarms.” *American Journal of Political Science* 28(1).

- McFadden, Daniel. 1974. Conditional Logit Analysis of Qualitative Choice Behavior. In *Frontiers in Econometrics*, ed. Paul Zarembka. New York: Academic Press pp. 105–142.
- Oberfield, Zachary W. 2014. *Becoming Bureaucrats: Socialization at the Front Lines of Government Service*. University of Pennsylvania Press.
- Ouss, Aurélie and John Rappaport. 2020. “Is Police Behavior Getting Worse? Data Selection and the Measurement of Policing Harms.” *Journal of Legal Studies* 49(1).
- Patty, John W. and Ian R. Turner. 2021. “Ex Post Review and Expert Policy Making: When Does Oversight Reduce Accountability?” *Journal of Politics* 83(1):23–39.
- Petersen, Nick, Ojmarrh Mitchell and Shi Yan. 2024. “Do progressive prosecutors increase crime? A quasi-experimental analysis of crime rates in the 100 largest counties, 2000–2020.” *Criminology & public policy* 23(2):459–490.
- Richman, Daniel. 2003. “Prosecutors and Their Agents, Agents and Their Prosecutors.” *Columbia Law Review* 103(4).
- Sasso, Greg. 2020. “Delegation and Political Turnover.” *Journal of Theoretical Politics* 32(2).
- Shjarback, John A, David C Pyrooz, Scott E Wolfe and Scott H Decker. 2017. “De-policing and crime in the wake of Ferguson: Racialized changes in the quantity and quality of policing among Missouri police departments.” *Journal of criminal justice* 50:42–52.
- Sklansky, David Alan. 2018. “The Problems with Prosecutors.” *Annual Review of Criminology* 1.
- Turner, Ian R. 2017. “Working Smart and Hard? Agency Effort, Judicial Review, and Policy Precision.” *Journal of Theoretical Politics* 29(1):69–96.
- Weaver, Vesla M and Amy E Lerman. 2010. “Political consequences of the carceral state.” *American Political Science Review* 104(4):817–833.

- Weingast, Barry R. and Mark J. Moran. 1983. "Bureaucratic Discretion or Congressional Control? Regulatory Policymaking by the Federal Trade Commission." *Journal of Political Economy* 91(5).
- Wood, B Dan and Richard W Waterman. 1991. "The dynamics of political control of the bureaucracy." *American Political Science Review* 85(3):801–828.
- Yogev, Dvir. 2025. "Holding justice accountable: Intensive vs. extensive margins in prosecutor elections." *Public Opinion Quarterly* 89(4):1087–1123.

## A The Charging-Standard Model

This appendix derives the reviewer choice, states the incumbent's dynamic problem, and considers an implementation-cost extension.

**Reviewer choice.** Given a candidate implemented standard  $s_t$ , reviewer  $a$  minimizes the loss in Equation (2). The first-order condition is

$$(\tau - \alpha_a) + \kappa(\tau - s_t) = 0.$$

Writing  $w = \kappa/(1 + \kappa)$  gives

$$\tau_{at}^0(s_t) = (1 - w)\alpha_a + ws_t,$$

which is Equation (5). The second derivative is  $2(1 + \kappa) > 0$ , so the best response is unique.

The source-specific extension replaces the baseline office standard  $s_t^*$  with  $\tilde{s}_{ait}^*$  from Equation (13). The equilibrium reviewer threshold becomes

$$\tau_{ait}^* = (1 - w)\alpha_a + w\tilde{s}_{ait}^*,$$

which is Equation (14). A common  $\kappa$  makes the response to  $\tilde{s}_t^*$  parallel. The source-specific interactions come from  $z_a B_t$  and  $h(o_i)E_t$ .

**Dynamic office choice and equilibrium.** Let  $s$  denote the inherited standard and  $Z$  the political state. For incumbent  $q$ , define the continuation value

$$\mathcal{C}_q(s'; s, Z) = \mathbb{E} [\mathbf{1}\{q' = q\}V_q(s', Z') + \mathbf{1}\{q' \neq q\}\bar{V}_q \mid s, Z, s'] .$$

Equation (4) can then be written as

$$V_q(s, Z) = \max_{s'} \{-\ell_q(s'; s, E) + \delta \mathcal{C}_q(s'; s, Z)\}.$$

For an interior choice, the first-order condition is

$$2(s' - \hat{s}_q(E)) + 2\zeta(s' - s) = \delta \frac{\partial \mathcal{C}_q(s'; s, Z)}{\partial s'}. \quad (20)$$

The left side contains the incumbent's current ideal point and the cost of changing the inherited standard. The right side is the value of how today's standard changes her political future.

At an election, let  $p_q(s', Z)$  denote the probability that the incumbent is retained. The continuation value includes

$$p_q(s', Z) \mathbb{E}[V_q(s', Z') \mid q' = q] + (1 - p_q(s', Z)) \bar{V}_q.$$

Its derivative contains  $\frac{\partial p_q(s', Z)}{\partial s'} (\mathbb{E}[V_q] - \bar{V}_q)$ . Because the incumbent values remaining in office, a standard that improves her retention probability changes her choice before control passes to a successor. Earlier choices can have the same effect when they change the political state inherited at the election. A successor begins with the standard chosen by the prior administration.

On compact state and action spaces with bounded continuous payoffs and a continuous transition, the Bellman operator has a fixed point. A unique maximizer yields the policy  $\sigma_q(s, Z)$  in Equation (6). Combined with the unique reviewer best response, these policies form the Markov-perfect Stackelberg equilibrium stated in Proposition 1.

**Linear-quadratic benchmark.** Hold the political regime and ideal point fixed, suppress the election transition, and write the incumbent's cost-to-go as  $J(s)$ . The dynamic problem

becomes

$$J(s) = \min_{s'} \{ (s' - \hat{s})^2 + \zeta(s' - s)^2 + \delta J(s') \}.$$

The solution has  $J(s) = P(s - \hat{s})^2 + C$ , where the positive value of  $P$  solves

$$\delta P^2 + (1 + \zeta - \delta\zeta)P - \zeta = 0.$$

The policy is

$$\sigma(s) = (1 - \rho)\hat{s} + \rho s, \quad \rho = \frac{\zeta}{1 + \zeta + \delta P} \in [0, 1).$$

The inherited standard therefore generates partial adjustment even when the principal's ideal point is constant. When the political state is expected to change, the continuation term in Equation (20) also moves the current standard. The adjustment can begin before an anticipated election.

**Implementation-cost extension.** An office may also bear an organizational cost when its standard pulls reviewers away from their preferred thresholds. Greater departures may require closer supervision or create resistance and retention costs. In the baseline model, let  $\lambda \geq 0$  and add

$$\lambda \int (\tau_a^0(s_t) - \alpha)^2 dF_t(\alpha)$$

to the period loss in Equation (3), where  $\tau_a^0(s_t) = (1 - w)\alpha + ws_t$ . The added cost equals  $\lambda w^2[(s_t - \bar{\alpha}_t)^2 + \sigma_{\alpha,t}^2]$ . It makes the dynamic policy depend on the active reviewers' mean preferred threshold. The reviewer best response and the mapping from thresholds to approval remain unchanged. Replacing  $s_t^*$  with the extended policy in Equations (9)–(10) therefore yields the same decomposition. The same result holds for the source-specific model after replacing  $s_t^*$  with  $\tilde{s}_t^*$  and holding the centered exposure terms fixed. A change in reviewer composition can now operate directly through the case weights and indirectly through the common standard. The decomposition assigns the latter response to common movement

because it changes how active reviewers decide. I set  $\lambda = 0$  in the baseline to preserve the clean separation between direct personnel reweighting and the dynamic common standard.

**Distribution of thresholds.** Hold the centered exposure terms fixed. Let reviewer types have mean  $\bar{\alpha}_t$  and variance  $\sigma_{\alpha,t}^2$ . The common part of the equilibrium discretionary threshold has mean  $(1-w)\bar{\alpha}_t + w\tilde{s}_t^*$  and variance  $(1-w)^2\sigma_{\alpha,t}^2$ . A change in  $\tilde{s}_t^*$  shifts every threshold by  $w\Delta\tilde{s}_t^*$  and leaves reviewers' baseline differences intact. A personnel change instead moves the case-weighted mean reviewer type. The empirical decomposition is first-order, so changes in higher moments of  $F_t$  do not receive a separate personnel contribution.

**Office-wide approval.** After holding case mix, policy, and the heterogeneous exposure terms fixed, the first-order approximation in Equation (10) gives

$$A_t \approx \mu - \bar{f}(1-w) \sum_{a \in \mathcal{R}_t} \pi_{at} \alpha_a - \bar{f} w s_t^*.$$

Subtracting the base-period expression gives

$$\Delta A_t \approx -\bar{f}(1-w) \sum_{a \in \mathcal{R}_0 \cup \mathcal{R}_t} (\pi_{at} - \pi_{a0}) \alpha_a - \bar{f} w (s_t^* - s_0^*),$$

which is Equation (11) after expressing both components on the approval scale. The personnel term changes when the case-weighted mean reviewer type changes. The common term changes when the implemented standard changes.

In the source-specific extension, replace  $s_t^*$  with  $\tilde{s}_t^*$  in both expressions and hold the centered exposure terms fixed. The personnel term is unchanged, while the common term becomes  $-\bar{f} w (\tilde{s}_t^* - \tilde{s}_0^*)$ .

**Comparative statics and identification.** The common threshold component is

$$w\tilde{s}_t^* = w[\sigma_{q_t}(s_{t-1}^*, Z_t) + \eta_B B_t].$$

Reputational pressure changes the incumbent’s current ideal point through  $E_t$ . Electoral pressure changes the continuation term in Equation (20), while an election may change  $q_t$ . Pressure from an anticipated principal can also enter directly through  $B_t$ . If the same event moves reputation and the incumbent’s retention prospects, the common path does not separate their effects. It also does not distinguish a shared response to an anticipated successor from an incumbent who adapts to remain in office. The heterogeneous terms in Equation (14) imply

$$\frac{\partial^2 \tau_{ait}^*}{\partial B_t \partial z_a} = w\lambda_B, \quad \frac{\partial^2 \tau_{ait}^*}{\partial E_t \partial h(o_i)} = w\lambda_E.$$

These cross-partials distinguish the sources only if  $z_a$  affects the response to hierarchy but not reputation and  $h(o_i)$  affects the response to reputation but not hierarchy. The exposure measures must also be predetermined with respect to the change in charging. A reviewer-specific  $\kappa_a$  would not supply this exclusion: it would make the reviewer more responsive to every source operating through  $s_t^*$ .

**Categorical rules and the estimating equation.** For cases with  $D_{it} = 0$ , the first-order approximation maps the threshold model to the estimating equation as

$$g(t) = -\bar{f}w\tilde{s}_t^*, \quad \tilde{\alpha}_a = -\bar{f}(1-w)\alpha_a = -\frac{\bar{f}\alpha_a}{1+\kappa},$$

with source-specific interaction coefficients  $-\bar{f}w\lambda_B$  and  $-\bar{f}w\lambda_E$ . The reviewer fixed effect in Equation (16) is therefore the approval-scale version of the structural private standard. A stricter private bar produces a lower reviewer effect, while a rise in  $\tilde{s}_t^*$  produces a decline in the office-wide component  $g(t)$ . The covariates  $x_i$  enter the linear probability model to approximate variation in the conditional distribution  $G(\cdot \mid x_i)$ . The policy coefficient  $\gamma$  does not equal a structural threshold parameter. The observed offense-by-post indicator is a proxy for  $D_{it}$ , and the retail policy was presumptive. Its coefficient combines the fraction of exposed cases that satisfy the eligibility conditions with compliance and any policy-induced

change in the cases reaching review. It is an intention-to-treat contrast.

## B Reviewer Assignment and Balance

The decomposition of the charging shift into personnel and incumbent-discretion channels relies on the claim that felony-review cases are assigned to on-call reviewers in a manner close to random. Two kinds of evidence support it. Institutionally, reviewers work fixed twelve-hour shifts and incoming calls are assigned in rotation to whoever is next in the shift’s order, so assignment depends only on the order of calls and the roster (Section 4; [Jordan 2024](#)). Statistically, case attributes should then explain little about which reviewer receives a case.

Table [A1](#) reports balance tests across the case attributes observed at review: charge severity, defendant age, defendant race, defendant sex, felony class, and broad offense category. Each row asks how much reviewer identity adds to a regression of a case attribute on time controls. Conditioning on month alone, reviewer identity adds about 1.2 percent of the variance in charge severity and about 0.8 percent for defendant age. Conditioning on month and offense type cuts these to 0.7 and 0.5 percent, and the incremental shares for the remaining attributes are of the same order. Because offense mix can vary across shifts and rosters, month-level regressions can register a reviewer–offense association even when cases are not assigned by offense; the offense controls therefore serve as a diagnostic rather than as part of the assignment rule. The relevant identifying comparison is among cases arriving under the same active roster at similar times, which the month controls only approximate, so I read these balance results as consistent with the institutional assignment account rather than as a reconstruction of the assignment blocks themselves. Consistent with quasi-random assignment, the between-reviewer share of the variance in approval, corrected for sampling noise with a split-sample estimator ([Kline, Saggio and Sølvsten 2020](#)), is 2.3 percent (plug-in: 3.2 percent). [Jordan \(2024\)](#) documents the rotation and reports complementary balance tests in an overlapping sample.

Table A1: Reviewer assignment balance: incremental variance in case attributes explained by reviewer identity.

| Case attribute        | Conditioning         | Base $R^2$ | + reviewer $R^2$ | Incremental |
|-----------------------|----------------------|------------|------------------|-------------|
| Charge severity       | month                | 0.008      | 0.020            | 0.012       |
| Charge severity       | month + offense type | 0.249      | 0.255            | 0.007       |
| Defendant age         | month                | 0.003      | 0.010            | 0.008       |
| Defendant age         | month + offense type | 0.093      | 0.098            | 0.005       |
| Black defendant       | month                | 0.005      | 0.018            | 0.012       |
| Black defendant       | month + offense type | 0.044      | 0.054            | 0.010       |
| Male defendant        | month                | 0.002      | 0.007            | 0.005       |
| Male defendant        | month + offense type | 0.022      | 0.027            | 0.004       |
| Class M/X or 1 charge | month                | 0.004      | 0.014            | 0.010       |
| Class M/X or 1 charge | month + offense type | 0.203      | 0.209            | 0.006       |
| Narcotics offense     | month                | 0.002      | 0.259            | 0.257       |

*Each cell reports the  $R^2$  from a linear regression of the case attribute on the conditioning set (fixed effects), with and without reviewer fixed effects.*

Because felony review is a rotating junior assignment, the reviewer panel is not a fixed cast, and the within-reviewer results should be read as movement within observed tenures rather than as following a large set of literal stayers across administrations. Table A2 reports the structure of the panel: the number of reviewers observed, the typical caseload per reviewer, the median observed tenure, the number of reviewers who handle cases on both sides of the December 2016 inauguration, and the share of all reviewed cases those spanning reviewers handle.

Table A2: Structure of the reviewer panel, 2014–2019.

| Quantity                                                  | Value |
|-----------------------------------------------------------|-------|
| Reviewers observed                                        | 550   |
| Reviewers with $\geq 25$ cases (estimation sample)        | 324   |
| Median cases per reviewer                                 | 56    |
| Median observed tenure (months)                           | 19    |
| Reviewers observed on both sides of the inauguration      | 224   |
| Share of all reviewed cases handled by spanning reviewers | 0.65  |
| Share of cases in 2016–2017 handled by spanning reviewers | 0.89  |

## C The Personnel Channel Under Panel Rotation

The estimate that personnel replacement carries a smaller share of the decline depends on how the composition channel is measured, and a natural alternative measure is misleading in this setting. A simple approach splits the pre-to-post change in approval into a between-reviewer component, from the changing share of cases each reviewer handles, and a within-reviewer component, from changes in each reviewer’s own rate. Applied to felony review, that split assigns about two-thirds of the decline to composition. That number is an artifact of the setting, not a finding.

The problem is panel rotation. Felony review is a junior posting, and most reviewers are observed on only one side of the December 2016 inauguration (Table A2). A between/within split has no pre-period rate for a reviewer who first appears after the transition, so it imputes one, and the natural imputation is the pre-period average. A large share of the reviewers active after the transition are such entrants (Table A2). Their approval rates lie on the same declining trend as everyone else’s, but because they have no measured pre-period baseline, the split books that trend-driven decline as “composition” rather than as movement within reviewers. The within component, by construction, can credit only reviewers observed on both sides of the transition. In an almost fully rotating panel the decomposition therefore loads the office-wide trend onto composition mechanically, whether or not the incoming reviewers differ from the outgoing ones.

The additive estimator in the body avoids this imputation. It estimates each reviewer’s stable fixed effect after removing the common change over time and case mix, then weights those effects by the reviewer’s actual case shares in 2014 and 2019. A reviewer who appears on only one side of the transition therefore receives a type estimate from that reviewer’s own decisions net of calendar time. The reviewer–month graph is fully connected, so overlap among active rosters links all 321 reviewers and all 72 months in the estimation sample. The point estimate assigns 18.7 percent of the adjusted decline to reviewer composition and 81.3 percent to the common component (Table 2). The 95 percent interval for the personnel

contribution includes zero, while the common component remains negative.

The result is stable across the specifications closest to the estimating equation. Allowing between six and twenty degrees of freedom in the common trend assigns between 17.0 and 18.4 percent of the decline to composition. Setting the minimum reviewer caseload between ten and one hundred cases produces estimates between 16.5 and 19.8 percent. Simpler diagnostics produce smaller estimates. A binary turnover measure attributes about 0.4 percent of the decline to reviewers first observed after the inauguration, while a leave-one-out approval propensity  $\psi$  attributes about four percent. Incoming reviewers are also slightly more approving than the incumbents they join. These diagnostics use different measures of personnel change, but none assigns most of the decline to replacement.

The direct within-reviewer comparison reaches the same qualitative conclusion. I restrict the comparison to 82 reviewers with at least ten decisions before the inauguration and ten after it, then estimate annual effects with reviewer fixed effects, the same case controls, and the retail-theft policy indicator used in the decomposition. Relative to 2016, adjusted approval in 2019 is 3.61 percentage points lower (95 percent confidence interval  $[-6.87, -0.35]$ ; Figure 3). Requiring between five and twenty decisions on each side produces declines between 3.45 and 3.90 points. Because felony review is a junior rotation, this evidence concerns movement within observed tenures rather than the conversion of a long-serving cast. It does not imply that every reviewer followed a parallel path. It also leaves open training, socialization, and the supervisors who translate an administration’s priorities into line practice (Bowman, Lowrey-Kinberg and Gould 2026, Oberfield 2014).

## D Pooled Retail-Theft Policy Specification

The pooled analysis uses felony-review decisions from January 1, 2014, through December 31, 2019. I begin with the same complete-case panel used for the office-wide analysis, requiring nonmissing values for the review outcome, review date, reviewer, defendant age, charge

severity, race, sex, and offense category. I retain reviewers who handled at least 25 cases in that panel before restricting the offense comparison. The main specification then limits the sample to cases classified as retail theft or ordinary theft. After removing fixed-effect singletons, the model contains 14,125 decisions by 312 reviewers across 72 calendar months.

Let  $R_i$  indicate retail theft and let  $P_{it} = \mathbf{1}\{d_i \geq \text{December 15, 2016}\}$ , where  $d_i$  is the review date. I estimate

$$\mathbb{E}[\text{approve}_{it} \mid a(i), t, x_{it}] = \mu_t + \rho R_i + \lambda P_{it} + \gamma(R_i \times P_{it}) + x'_{it}\beta + \tilde{\alpha}_{a(i)}, \quad (21)$$

where  $\mu_t$  is a calendar-month fixed effect and  $\tilde{\alpha}_{a(i)}$  is a reviewer fixed effect. The vector  $x_{it}$  contains linear controls for defendant age and charge severity and fixed effects for race and sex. The retail-theft indicator  $R_i$  captures the baseline difference between the two offenses. Calendar-month effects absorb changes shared by the offenses across months. Because the policy begins halfway through December 2016,  $P_{it}$  also varies within that month; including its main effect absorbs any change after December 15 shared by retail and ordinary theft. The interaction coefficient  $\gamma$  is the differential post-policy change in retail-theft approval. Standard errors are clustered by calendar month.

The main specification uses the policy's effective date. I estimate three alternatives. First, I begin the post-policy period in January 2017, the first full month after the rule. The post indicator is then collinear with the calendar-month effects, so that model includes only its interaction with retail theft. Second, I replace ordinary theft with a broader property-offense comparison group: theft, burglary, residential burglary, possession of a stolen motor vehicle, forgery, identity theft, credit-card cases, criminal damage to property, fraud, deceptive practice, theft by deception, and aggravated identity theft. Third, I use all offense categories except driving with a suspended or revoked license, which was subject to a separate policy during the study period. The broader comparison models retain the December 15 cutoff and the common post-date main effect.

Table A3 reports the results. Relative to ordinary theft, retail-theft approval falls by 3.56 percentage points after December 15 (95 percent confidence interval  $[-5.77, -1.34]$ ). Beginning the post period in January 2017 produces a 3.26-point decline ( $[-5.53, -1.00]$ ). The estimate is 4.08 points using other property offenses as the comparison ( $[-5.38, -2.77]$ ) and 5.04 points using all offenses except suspended-license cases ( $[-6.22, -3.85]$ ). The result is therefore not specific to the transition-month coding or to ordinary theft as the comparison. Because the data do not record every eligibility condition and the policy changes which retail-theft cases reach review, these coefficients remain offense-level reduced-form estimates rather than structural effects of  $D_{it}$ .

Table A3: Pooled estimates of the retail-theft policy response.

| Comparison                               | Difference (p.p.) | 95% CI           | Cases  |
|------------------------------------------|-------------------|------------------|--------|
| Ordinary theft; Dec. 15 cutoff           | -3.56             | $[-5.77, -1.34]$ | 14,125 |
| Ordinary theft; Jan. 2017 cutoff         | -3.26             | $[-5.53, -1.00]$ | 14,125 |
| Property offenses; Dec. 15 cutoff        | -4.08             | $[-5.38, -2.77]$ | 30,923 |
| All offenses except DWLS; Dec. 15 cutoff | -5.04             | $[-6.22, -3.85]$ | 88,997 |

*Each row reports the coefficient on retail theft interacted with the indicated post-policy period. All models include reviewer and calendar-month fixed effects, linear controls for age and charge severity, and fixed effects for race and sex. Models using the December 15 cutoff also include the post-date main effect. Standard errors cluster by calendar month. The estimation sample and comparison groups are described in the text.*

## E Timing Tests

The main timing specification uses the full period because the model permits gradual adjustment and the political events are close together. Table A4 shows how the joint level-and-slope test changes as the window narrows. The McDonald-video result persists in every window. The primary and inauguration show no joint change in the eighteen-month, twenty-four-month, or full-period specifications. Their twelve-month estimates are unstable because those windows include adjacent events and ask a small number of monthly observations to estimate two trends and a level shift.

Panel B reports minimum detectable effects with 80 percent power, in percentage points. The full-period specification can detect immediate shifts of 0.50 points at the video, 0.80 at the primary, and 0.91 at inauguration. It is less informative about gradual changes: the minimum detectable one-year change is 1.32 points at the primary and 1.24 points at inauguration. Ninety-percent equivalence intervals fall within a one-point bound for both immediate electoral-date shifts. The annual inauguration slope and its combined one-year effect do not meet that equivalence standard.

The placebo grid estimates the full-period segmented model at each of 48 interior months. The video ranks third by the joint level-and-slope statistic. The primary ranks forty-seventh and inauguration twenty-fifth. The grid confirms that the electoral dates do not stand out as unusual breaks, while the end of 2015 does.

Table A4: Sensitivity and detection limits for candidate-date timing tests.

| Candidate date                                                | Estimation window around candidate date |                  |              |              |
|---------------------------------------------------------------|-----------------------------------------|------------------|--------------|--------------|
|                                                               | 12 months                               | 18 months        | 24 months    | Full period  |
| <i>Panel A: Joint p-value for level and slope change</i>      |                                         |                  |              |              |
| McDonald video                                                | < .001                                  | < .001           | < .001       | 0.001        |
| Democratic primary                                            | 0.001                                   | 0.122            | 0.700        | 0.553        |
| Inauguration                                                  | < .001                                  | 0.292            | 0.203        | 0.162        |
| <i>Panel B: Full-period detection limits and placebo rank</i> |                                         |                  |              |              |
| Candidate date                                                | Level MDE                               | Annual slope MDE | One-year MDE | Placebo rank |
| McDonald video                                                | 0.50                                    | 0.86             | 0.96         | 3/48         |
| Democratic primary                                            | 0.80                                    | 1.05             | 1.32         | 47/48        |
| Inauguration                                                  | 0.91                                    | 1.00             | 1.24         | 25/48        |

*Panel A reports the joint p-value for the immediate level and post-date slope terms. Panel B reports 80-percent minimum detectable effects in percentage points and the descending rank of each candidate date's joint statistic among 48 admissible dates. All specifications are case-level linear probability models with reviewer fixed effects, case controls, and the retail-theft policy indicator. Newey–West standard errors aggregate scores by month and use a six-month lag.*

## F Selection into Review and Bounds on the Within-Pool Change

The population entering felony review changed substantially over the window, and this appendix establishes that the change cannot account for the decline in approval. Chicago felony arrests fell from about twenty-five thousand in 2014 to about eighteen thousand in 2016–2017 and then recovered to about twenty-two thousand by 2019 (Figure A2). The largest decline was in narcotics-only arrests, which bypass felony review. Among review-eligible arrests, the share reaching review fell from about 0.97 in 2014 to about 0.90 in 2019 (Figure A3; Table A5). The offense and racial composition of the reviewed pool also changed. Aggregate charging rates therefore combine movement in the office’s standard with movement in the cases presented to it.

The retail-theft policy produces a concentrated version of the same selection problem. Retail theft accounts for about fifteen percent of reviewed cases before the rule and about five percent afterward (Figure A1). The fall occurs immediately and persists. Approval among the retail-theft cases that remain in review therefore conditions on a sharply altered case pool.

I bound the within-pool change with trimming bounds in the style of Lee (2009). The bounds are not assumption-free. They rest on one assumption about selection: monotonicity. When the selection rate falls from  $s_0$  to  $s_1 < s_0$ , the tighter regime reviews a subset of the cases the looser regime would have reviewed; screening removes cases at the margin and adds none. Monotonicity defines the population the bounds cover: the cases that reach review under either regime, which is exactly the tighter period’s reviewed pool. The looser period’s pool contains that population plus a marginal fraction  $q = (s_0 - s_1)/s_0$ , and the data do not reveal where those marginal cases sit in the approval distribution. Deleting the bottom  $q$  of the looser period’s distribution yields the highest approval rate for the common population consistent with the data; deleting the top  $q$  yields the lowest. The result is an interval

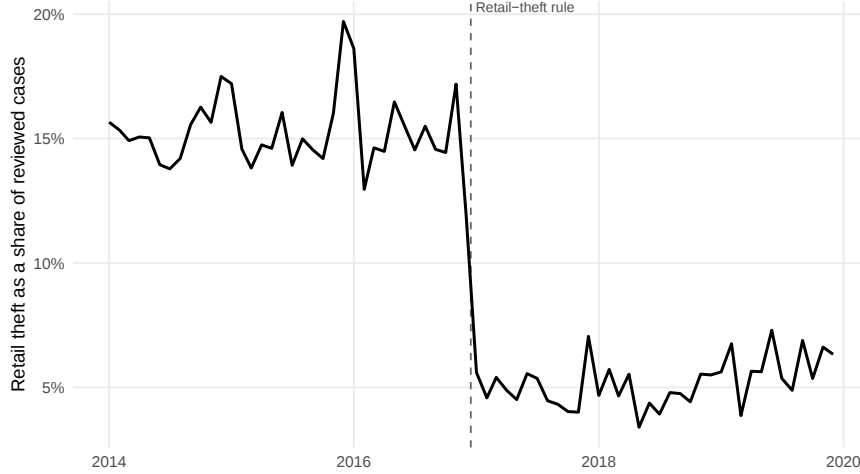


Figure A1: Retail theft as a share of the felony-review docket. The share falls from about 0.15 to about 0.05 when the December 2016 declination policy takes effect and remains near that level through 2019.

$\Delta_t^L \leq \Delta_t \leq \Delta_t^U$  for the within-pool change, worst-case with respect to how the marginal cases sort on the outcome. Monotonicity is what separates these bounds from bounds that impute extreme outcomes for every unreviewed case (Manski 1990). I report the worst-case bounds and bounds tightened by the same trimming within race-by-offense-class strata.

The bounds indicate that selection cannot account for the decline in approval. From 2016 forward, the worst-case interval for the within-review change lies below zero. Trimming within race-by-offense-class strata narrows the bounds further; for 2018, the covariate-tightened interval is approximately  $[-0.036, -0.027]$  (Figure A4; Table A5). Inference does not overturn the sign: the Imbens–Manski (Imbens and Manski 2004) 95 percent confidence intervals on the bounds remain below zero from 2016 forward, with the 2018 worst-case interval at about  $[-0.044, -0.026]$ . Selection also moves in the opposite direction from the observed decline. Entry into review becomes more restrictive over the window. If the cases retained by stricter approval standards are positively selected on strength, the observed decline in approval understates the change in the office’s standard. The office-wide movement is therefore not produced by a weaker or more permissively selected pool.

I read these bounds as a sensitivity analysis and check the reading three further ways.

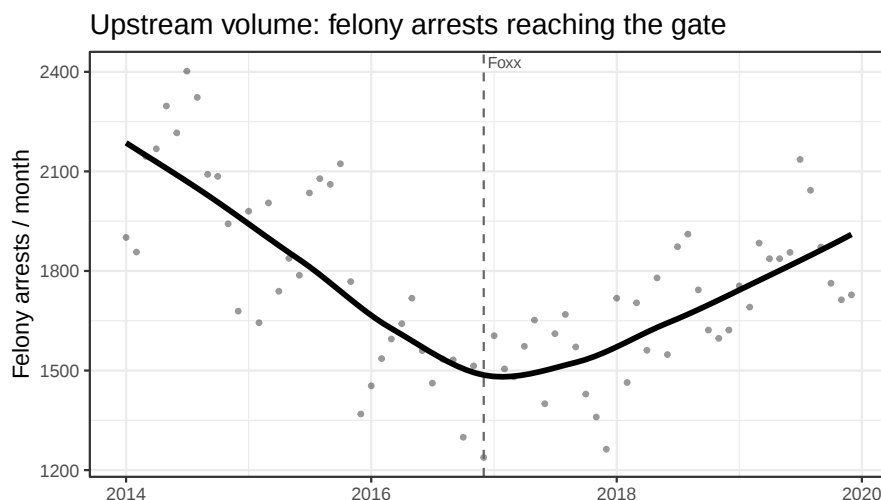


Figure A2: Chicago felony arrests fall by roughly a quarter into 2016–2017 before partially recovering toward 2019. The series is all felony arrests, including the narcotics-only arrests that bypass review.

First, reweighting each year’s reviewed pool to the 2014 distribution of race and felony class leaves approval almost unchanged, within a fraction of a percentage point of the raw path, so drift in the composition of the pool does not produce the decline. Second, coding every unlinked review-eligible arrest in the most adverse way, approved in the base year and declined thereafter, deepens the decline rather than erasing it, so unmatched bookings cannot manufacture it. Third, dropping the last quarter of 2019, where the arrest-to-case match rate softens, leaves the trajectory intact. The selection account survives each.

## G Linkage Diagnostics

This appendix documents the arrest-to-case linkage on which the funnel depends and shows that it is stable enough to read movements in the selection rate as the selection process itself. The link is built from a FOIA’d arrest-information file that carries both the Chicago Police booking number and the State’s Attorney’s case-participant identifier (Section 4); the booking number joins to the arrest universe and the case-participant identifier to the charging and disposition records. About ninety-five percent of felony arrests match to a case

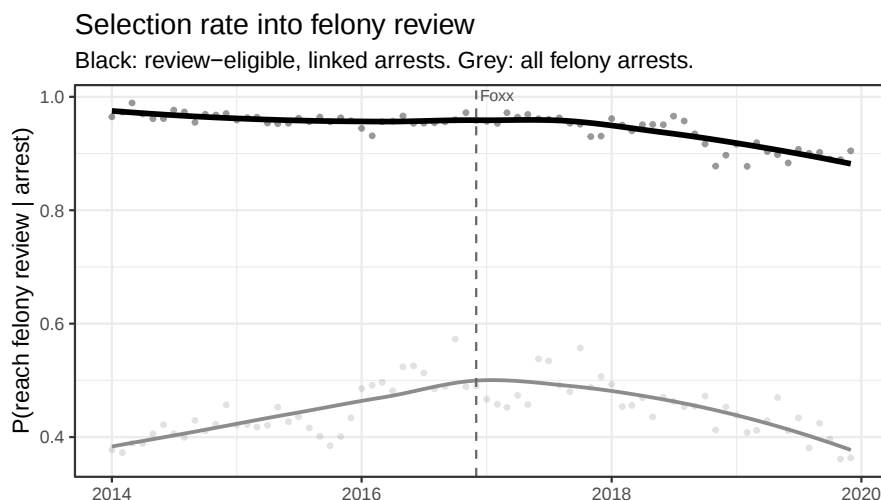


Figure A3: Selection into felony review. Black: among review-eligible arrests, the share reaching felony review declines from about 0.97 in 2014 to about 0.90 in 2019. Grey: all felony arrests, where the level is dominated by narcotics-only arrests that bypass review by rule and by unlinked bookings. The pool entering the funnel is endogenous to the enforcement environment.

record.

The concern is not the level of the match rate but its stability: a link whose coverage drifted over time could manufacture the very selection patterns the paper studies. Figure A5 shows it does not. The monthly match rate is flat through 2018 and softens only in the last quarter of 2019. Two features of the design keep even that late softening out of the results. First, the selection rate is computed conditional on a successful link, so a change in match coverage does not enter the measure mechanically. Second, the analyses that use the arrest link are robust to the drift: dropping the last quarter of 2019 leaves the trajectory intact, and coding every unlinked review-eligible arrest in the most adverse way—approved in the base year, declined thereafter—deepens the estimated decline rather than erasing it (Appendix F). The roughly five percent of arrests that never link blend genuine non-prosecution with imperfect booking numbers, which the data cannot fully separate; the robustness checks bound the consequences of that ambiguity rather than resolve it.

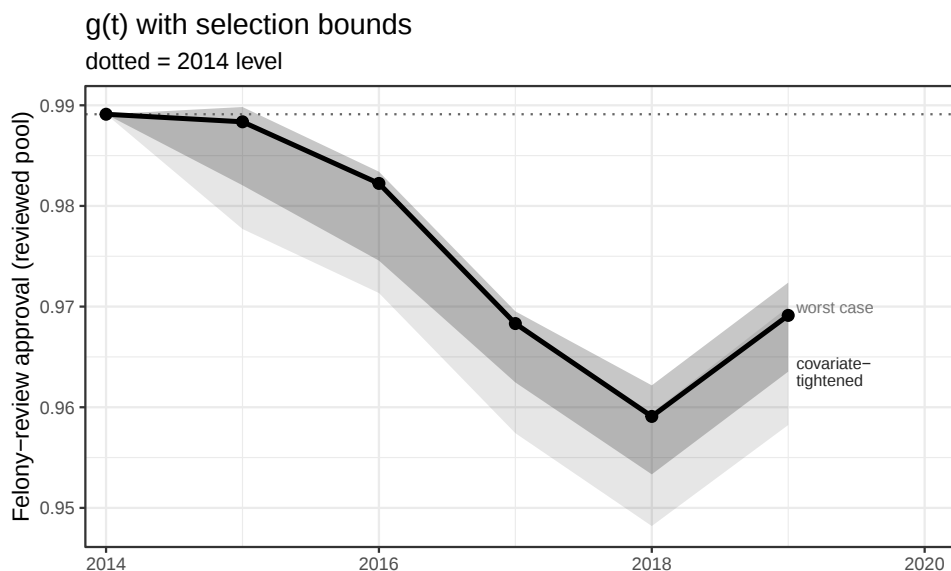


Figure A4: Trimming bounds on the within-review change in approval, computed on the review-eligible universe. Even the worst-case bounds (light) place the 2016–2019 changes below zero; the covariate-tightened bounds (dark) narrow them further. Both sets impose monotone selection.

## H The Structure of Screening

If the office had begun screening more selectively on evidence, the approve-or-decline decision should have become more responsive to case characteristics, and the caseload should have tilted toward offenses in which investigation reliably yields evidence. Neither occurs. The share of the approval decision explained by case characteristics falls over the window (McFadden pseudo- $R^2$  of 0.098 in 2014 and 0.031 in 2019; [McFadden 1974](#)), and the evidence-productive share of the reviewed caseload is nearly flat (0.55 in 2014, 0.58 in 2019). The change is a broadly uniform reduction in the level of approval, not a new screening rule.

## I Downstream Outcomes

These analyses describe disposition outcomes among defendants whose cases were approved for felony prosecution. Because approval is endogenous to the office’s charging standard, they characterize what changed over the same period rather than identify an effect of the

Table A5: Selection into review and bounds on the within-review change in approval, relative to 2014.

| Year | Selection rate $S_t$ | Raw $\Delta$ approval | Worst-case [lo, hi] | Covariate-tightened | Imbens–Manski 95% CI |
|------|----------------------|-----------------------|---------------------|---------------------|----------------------|
| 2015 | 0.959                | −0.001                | [−0.011, −0.001]    | [−0.007, +0.001]    | [−0.015, +0.002]     |
| 2016 | 0.955                | −0.007                | [−0.018, −0.007]    | [−0.015, −0.006]    | [−0.020, −0.004]     |
| 2017 | 0.956                | −0.021                | [−0.032, −0.021]    | [−0.027, −0.020]    | [−0.035, −0.017]     |
| 2018 | 0.938                | −0.030                | [−0.041, −0.030]    | [−0.036, −0.027]    | [−0.044, −0.026]     |
| 2019 | 0.900                | −0.020                | [−0.031, −0.019]    | [−0.026, −0.017]    | [−0.034, −0.016]     |

$S_t$  is the annual share of review-eligible felony arrests (narcotics-only arrests, which bypass review by rule, are excluded) that reach felony review, conditioned on a successful arrest-to-case link; the 2014 baseline is 0.969. Because the selected share is high, even the worst-case bounds sign the decline from 2016 on. Both bound sets impose monotone selection; the worst-case columns impose nothing else. The covariate-tightened bounds trim within race  $\times$  offense-class strata. The final column reports the Imbens–Manski 95 percent confidence interval for the worst-case bound. From 2014 to 2019 the selection rate fell by 0.069 while approval fell by 0.020, so positive selection implies the observed decline understates the within-pool shift.

changing standard.

**Disposition data.** The State’s Attorney’s disposition records report the resolution of each charge. I link them to the initiation records through the case-participant identifier. The disposition data run through late 2024, so I observe the final resolution of every case entering the 2014–19 sample.

**Conviction.** I measure conviction two ways. Conviction on the primary charge records whether the defendant is found or pleads guilty on the leading charge. Any-charge conviction records whether the defendant is convicted on any charge in the case. The latter measure accounts for secondary charges that are filed and later dropped, a practice these same data show is applied unevenly across defendants (Clark forthcoming). I separately record whether a conviction was secured by guilty plea.<sup>16</sup>

<sup>16</sup>A charge counts as a conviction if its disposition is one of: Plea Of Guilty; Plea of Guilty - Amended Charge; Plea of Guilty - Lesser Included; Finding Guilty; Finding Guilty - Amended Charge; Finding Guilty - Lesser Included; Verdict Guilty; Verdict Guilty - Amended Charge; Verdict Guilty - Lesser Included. The list excludes dispositions that contain the word “guilty” but are not convictions (“Verdict-Not Guilty,” “Finding Not Not Guilty” [sic]). Guilty pleas are the three plea categories.

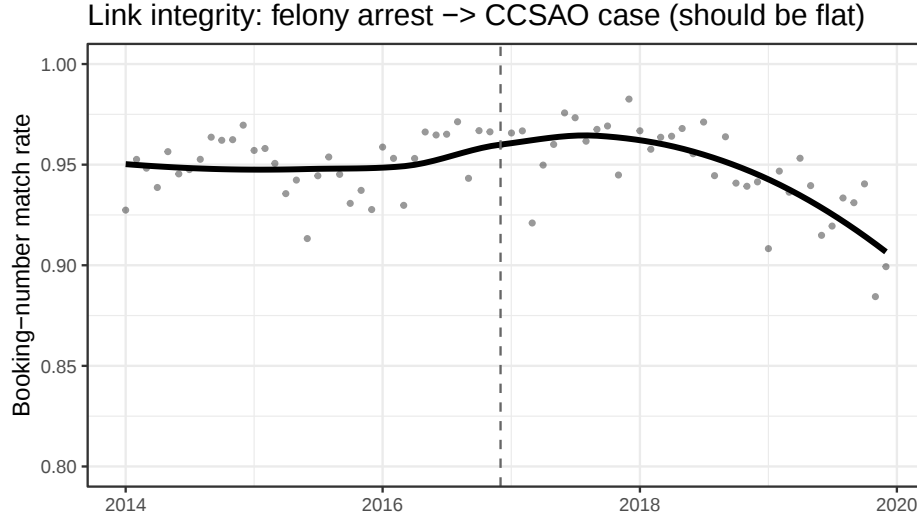


Figure A5: Booking-number match rate by month. A link that drifted over time could manufacture selection patterns; it is flat through 2018 and softens only in late 2019. The selection rate conditions on a successful link, which removes this drift from the measure.

**Conviction severity.** I use the offense-class ordering defined in the main text to compare the charge filed at initiation with the charge of conviction. I code a case as *downgraded* when the convicted class falls below the charged class. This measure captures the reduction in severity between what the office charges and what ultimately sticks.

**Reason for dismissal.** For cases that end without conviction, I classify the recorded disposition reason as evidentiary (suppression, laboratory, or witness problems), discretionary (diversion-program completion or a declination on the merits), or administrative (the case proceeds on another count or by another route). The classification distinguishes cases that failed from cases the office chose not to pursue.

Leniency can operate on two margins (Yogev 2025). An office can withdraw from cases it once would have carried to conviction, or it can continue to obtain convictions while reducing their severity. Reductions from felony to misdemeanor charges can have important consequences for defendants (e.g., Agan, Doleac and Harvey 2023, Jordan 2025). In Cook County, the second margin changed more. Conviction on some charge among approved defendants declined from about 0.85 in 2014 to about 0.80 in 2019, while guilty pleas on

some charge remained at about 0.79. The guilty plea to the original lead charge fell from about 0.67 to about 0.60. More sharply, the share of approved defendants convicted of a lower offense class than the one charged rose from about 0.19 to about 0.43. Among convictions, the felony share fell from about 0.92 to about 0.81 (Table A6; Figure A6). Both changes survive adjustment for offense mix. The downgrade rate uses approved defendants as the denominator and counts the unconvicted as zeros, so the increase is not produced by a shrinking base of convictions.

Table A6: Downstream outcomes among approved cases: conviction remains common while severity falls.

| Outcome (among approved)    | Pre (2014) | Post (2019) | Offense-adj. $\Delta$ |
|-----------------------------|------------|-------------|-----------------------|
| Conviction, any charge      | 0.85       | 0.80        | —                     |
| Conviction, lead charge     | 0.76       | 0.60        | −0.10*** (0.006)      |
| Guilty plea, any charge     | 0.79       | 0.78        | +0.006 (0.006)        |
| Acquittal                   | 0.047      | 0.021       | −0.042*** (0.002)     |
| Dismissal                   | 0.10       | 0.175       | +0.084*** (0.004)     |
| Charged→convicted downgrade | 0.19       | 0.43        | +0.18*** (0.012)      |
| Felony share of convictions | 0.92       | 0.81        | −0.10*** (0.004)      |

*Offense-adjusted  $\Delta$  is the coefficient on review-month  $\times 60$  ( $\approx$  full window) from a linear model with offense fixed effects, with standard errors (in parentheses) clustered by review month; \*\*\* $p < 0.001$ .*

*Any-charge conviction is the convictability margin, shown descriptively. Guilty plea here is measured on any charge. The downgrade rate is measured among approved defendants, with the unconvicted counted as zero; the felony share is measured among convictions.*

Dismissal also rose, from about 0.10 to about 0.175, so some leniency occurred through withdrawal. Among dismissals with a recorded reason, the increase is in discretionary rather than evidentiary drops. Blank or missing reasons are the largest rising category, however, so the composition of that increase is only suggestive. Acquittals fall from about 0.047 to about 0.021, and the evidentiary dismissal rate per approved case is flat to slightly lower even as the overall dismissal rate rises (Figures A8 and A9). The cases carried forward are therefore less likely to lose at trial or collapse for evidentiary reasons. These patterns are consistent with selection at a higher charging bar, but do not provide a separate test of why the standard changed.

The guilty plea remains about 78 percent of approved cases throughout. The office con-

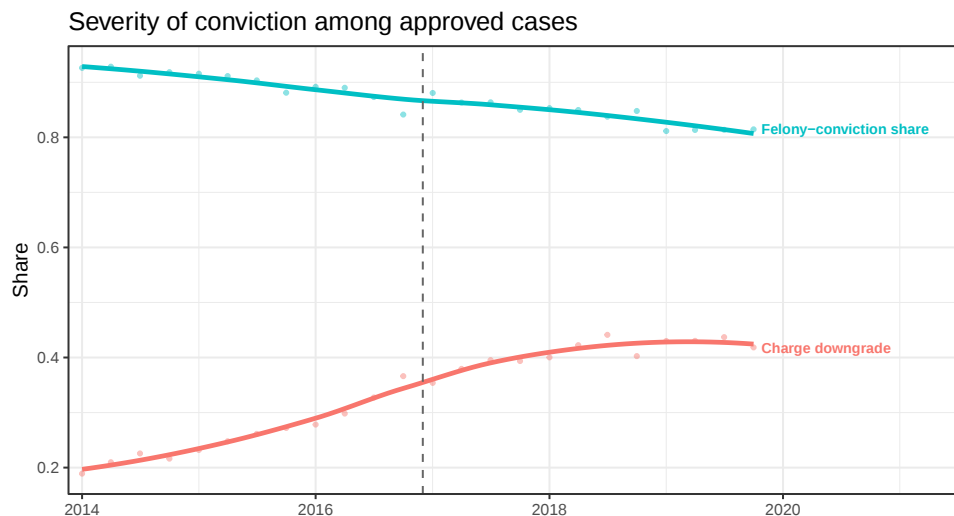


Figure A6: Charge reduction over the window, among approved defendants: the share convicted of a lower class than charged roughly doubles ( $\approx 0.19 \rightarrow 0.43$ ); among convictions, the felony share falls ( $\approx 0.92 \rightarrow 0.81$ ). Conviction severity falls more than conviction frequency.

tinued to rely on the same plea-based process while reducing the severity of the resulting conviction. The pattern is consistent with a system in which charge bargaining provides an easier route to moderation than the complete withdrawal of a case (see [Patty and Turner 2021](#), [Turner 2017](#), on court oversight and agency behavior); the formal argument is developed in [Clark \(2026b\)](#).

## Appendix References

Jordan, Andrew. 2024. “Racial Patterns in Approval of Felony Charges.” Working paper, Washington University in St. Louis.

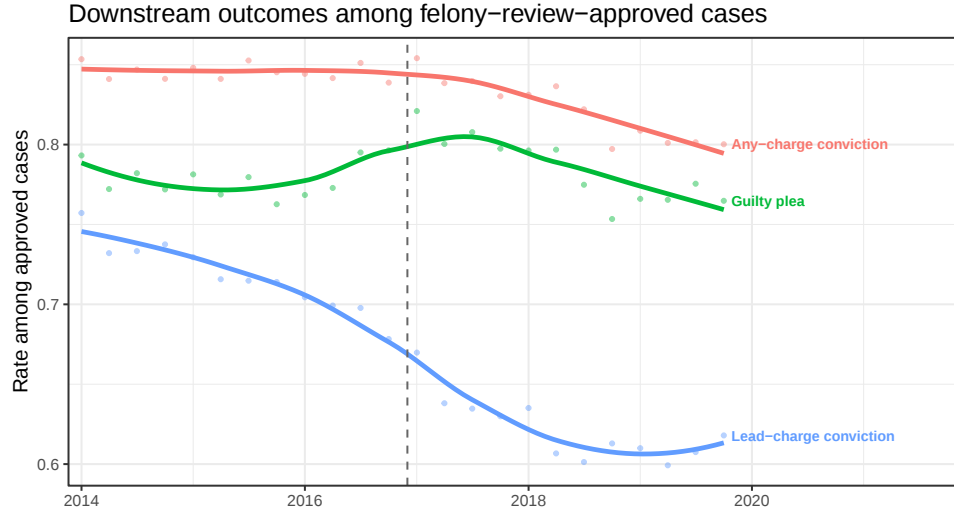


Figure A7: Conviction on some charge among approved defendants declines modestly across review cohorts ( $\approx 0.85 \rightarrow 0.80$ ) as approval falls. The downstream case pool does not show a comparable loss of convictability.

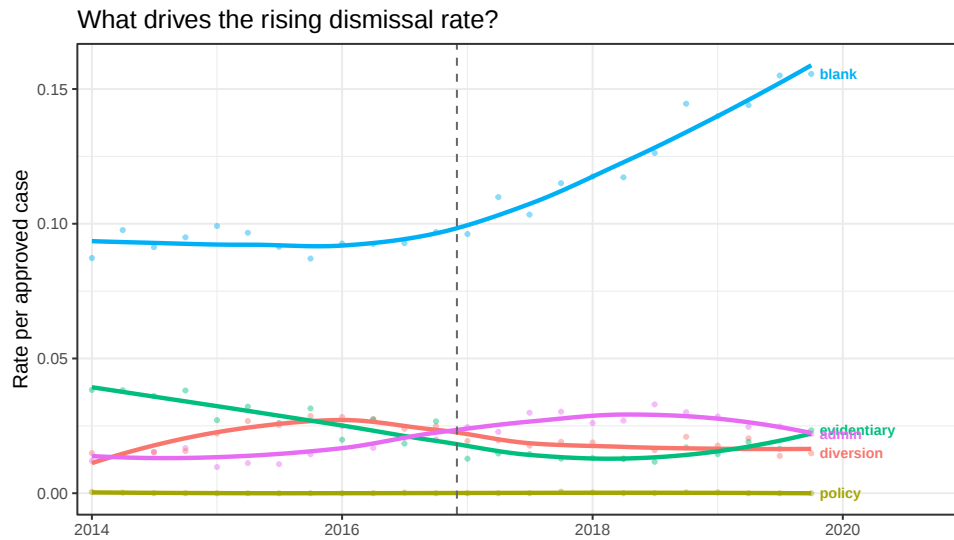


Figure A8: Dismissals of approved cases decomposed by recorded reason, as a rate per approved case (the reason-specific rates sum to the dismissal rate). As the overall dismissal rate rises, the evidentiary component is flat to slightly lower, so the additional dismissals are not cases that fell apart on the evidence; the largest rising component is blank or missing reasons.

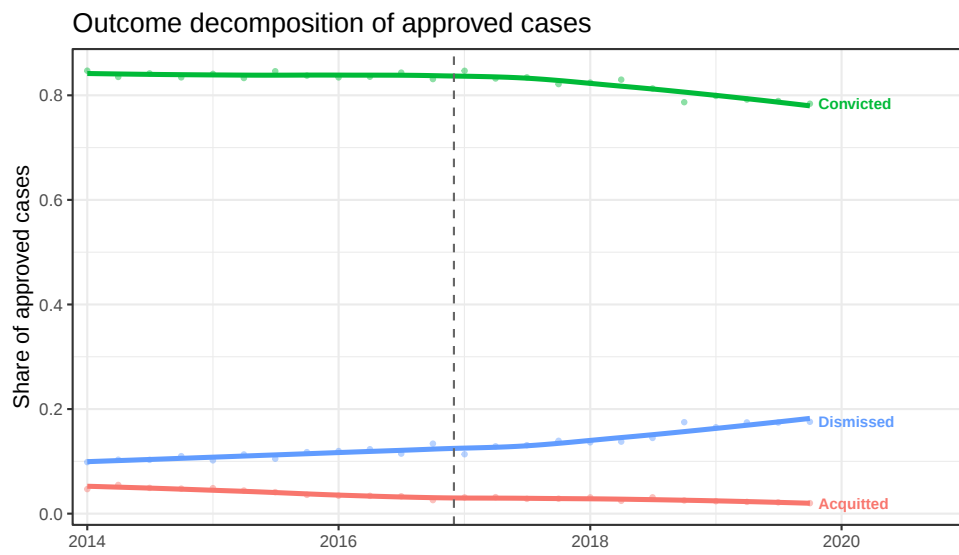


Figure A9: Disposition of approved cases over the window. Conviction remains the dominant outcome; dismissal rises modestly ( $\approx 0.10 \rightarrow 0.175$ ), concentrated in discretionary drops, while acquittal falls.