

Policy Advice and Electoral Accountability

Francis William Meda

Abstract

A common concern about expert advice in politics is that it can undermine accountability: by claiming to have "followed the expert," politicians can pursue private agendas without facing electoral punishment. I identify when this concern is justified by analyzing a model in which an incumbent either sets policy directly or first consults an expert advisor whose recommendation is private. I show that when voters cannot distinguish genuine, policy-improving consultation from consultation used to mask rent-seeking, the expert grants incumbents the cover they lack when implementing policy without advice. This cover arises only when the advisor is informative enough that some incumbents genuinely defer to her, and where it arises, accountability weakens on two margins: more incumbents pursue policies that are unlikely to serve voters, and those who do so are punished less. Therefore, the possibility of seeking experts' advice can reduce voter welfare, despite the genuine informational gains it provides.

1 Introduction

Expert advice has become a central feature of contemporary policymaking. Government officials have relied upon expertise in areas such as monetary and fiscal policy, national security, public health, and environmental regulation to justify transformative decisions and restrictions on rights. When the United Kingdom imposed its first national lockdown in March 2020, the government justified the decision by repeatedly invoking the advice of its Scientific Advisory Group for Emergencies, and “following the science” became the governing refrain (Ehlers and Esselborn, 2022). Subsequent inquiries revealed that the advisory body had, at several points, urged earlier or stricter measures than the government adopted (McKee, Hanson and Abbasi, 2022). While scholars acknowledge the substantive benefits of expertise in public decision-making, critics have long emphasized its symbolic and political uses: legitimating contested choices and deflecting responsibility for unpopular ones (Weiss, 1979; Amara, Ouimet and Landry, 2004; Boswell, 2009; Blum and Straßheim, 2024). They fear that officials who invoke expertise can escape responsibility for a decision voters would otherwise punish.

Despite these concerns about how expert advice affects officials’ accountability to voters, most formal theoretical work on expertise focuses on the agency relationship between politicians and their advisors. These models consider settings in which an advisor is better informed than the official she serves but may not share his objectives, and they ask when informative advice can be transmitted at all (Crawford and Sobel, 1982; Gilligan and Krehbiel, 1987; Calvert, 1985; Krishna and Morgan, 2001; Denisenko, Hafer and Landa, 2026). Officials in these models can, in principle, make better decisions by consulting, but they face the risk that an advisor with divergent preferences will shade her advice, so the quality and frequency of advice depend on how close the two are. While providing many insights about when advice is informative and how advisory institutions should be designed, this literature is built around a sender and a receiver with no audience. It is therefore largely silent on how electoral motivations shape the decision to consult, and on the extent to which consultation allows officials to escape accountability for policies that are not in their constituents’ interests. This paper addresses this limitation by developing a model of expert advice with electoral accountability. The model builds on existing models of political agency in which voters use observed actions to infer an incumbent’s type and reward those believed to share their interests (Ferejohn, 1986; Maskin and Tirole, 2004; Fox and Jordan, 2011).¹

The model features a representative voter and an incumbent politician. The voter has state-contingent policy preferences (e.g., if the policy question is whether to impose a lockdown

¹See Ashworth (2012) for a review.

during an epidemic, she wants the government to impose a lockdown only if the rate of contamination is high). The voter's problem is that she does not know the underlying state of the world (e.g., whether the rate of contamination is high or low). The politician must either set policy directly or first consult an expert advisor who knows more about the state than he does. If consulted, the advisor observes an informative signal of the state and privately recommends the policy the signal favors.

In addition to uncertainty about the state, the voter is uncertain about the politician's preferences. Some incumbents share her state-contingent preferences (congruent); others are non-congruent, with state-independent preferences (e.g., a politician who does not want to impose lockdown regardless of the true severity of the outbreak, say, to appease donors that benefit from keeping the economy open). Ex ante, a non-congruent politician's preferred policy is less likely to serve the voter's interest (unfavorable policy); that is, it is less likely to match the unknown state. The voter observes whether the incumbent consulted the expert and which policy he chose, but not what the expert recommended. After observing the consultation's decision and the policy implemented, the voter updates her belief that the incumbent shares her policy preferences, and the incumbent's reelection probability increases with that belief. I measure accountability on two dimensions. The first is an extensive margin (deterrence): the proportion of non-congruent incumbents who are deterred from implementing the unfavorable policy. The second is an intensive margin (punishment): the probability of reelection after implementing the unfavorable policy.

The benchmark model without advice captures ordinary electoral discipline. A congruent politician always implements the policy that is, ex ante, more likely to serve the voter's interest when acting without additional information. A non-congruent incumbent, therefore, faces a trade-off: he can implement the unfavorable policy and reveal himself as non-congruent, or he can mimic the congruent type and improve his electoral prospects. In this benchmark, accountability in the intensive margin operates because the voter can attach a clear reputational meaning to the unfavorable policy. It is chosen only by incumbents whose preferences diverge from hers, and so it is punished accordingly. As a result, few non-congruent incumbents implement the unfavorable policy.

The possibility of seeking private advice can change this inference. The reason is that there exist a compliance equilibria in which a congruent politician seeks advice and complies with the advisor's recommendation with a positive probability and a non-congruent politician seeks advice and always implements his preferred policy. Since the advisor is informative, genuine advice-following sometimes leads the congruent incumbent to choose the ex ante unfavorable policy. However, this same observable behavior is also produced by a non-congruent incumbent who consults and always chooses his preferred policy. The voter observes consultation and

the final policy, but not the recommendation. She therefore cannot tell whether consultation followed by an unfavorable policy reflects deference to expert information (consultation that improves policy: instrumental use of expertise) or the strategic use of consultation as cover (consultation that disguises rent-seeking: symbolic use of expertise).

This observational equivalence is what allows private advice to undermine accountability. In a compliance equilibrium, both types of politicians consult, but for different reasons. The congruent incumbent uses advice instrumentally to improve the match between policy and the state of the world. The non-congruent incumbent uses advice symbolically to make his preferred policy less revealing of bias. As a result, a non-congruent incumbent who pursues a policy contrary to the voter's *ex ante* interest after consulting is not punished as harshly as he would be if he had chosen the same policy without advice. Moreover, because the electoral cost of choosing that policy falls, more non-congruent incumbents choose it. Therefore, in a compliance equilibrium, advice weakens accountability on both margins: it reduces deterrence and softens punishment.

A compliance equilibrium exists when a congruent politician's policy benefit of following the advice exceeds the electoral cost of sometimes implementing the *ex ante* unfavorable policy. Thus, accountability failure is not driven by bad advice. It arises precisely because the advisor is useful and because congruent incumbents sometimes follow her into actions that voters would otherwise interpret as suspicious. The symbolic value of consultation is, therefore, parasitic on its instrumental value: non-congruent incumbents can borrow credibility from the fact that congruent incumbents genuinely defer to expertise. If the advisor were uninformative, the congruent incumbent would have little reason to bear the reputational cost of following advice into an unfavorable policy, and the equilibrium would reproduce the benchmark. If the advisor is sufficiently informative, however, the compliance equilibrium is the unique refined equilibrium, and accountability is necessarily undermined.

The welfare implications are ambiguous. In a compliance equilibrium, voters gain because congruent incumbents replace prior-based decisions with more accurate advice. But voters lose because some non-congruent incumbents who would have been disciplined without advice now choose their privately preferred policy under advisory cover. Advice improves welfare only when the informational gain for congruent incumbents outweighs the additional shirking by non-congruent ones. Thus, accountability failures and welfare improvements can coexist: expert advice may make voters better off even while weakening electoral discipline, and it may make them worse off when the cover it creates is sufficiently costly.

The paper proceeds as follows. Section 2 reviews the literature. Section 3 presents the model and the equilibrium refinement. Section 4 analyzes the no-advice benchmark,

characterizes the equilibria of the advice game that survive the refinement, and states the accountability results. Section 5 compares voter welfare with and without private advice. Section 6 discusses the mechanism, the scope conditions, and the natural disclosure extension and concludes.

2 Related Literature

This paper is primarily a contribution to the formal study of policy advice. Canonical models take a privately informed advisor and an uninformed decision-maker and ask when informative communication is possible at all, with the central finding that divergence in preferences limits revelation ([Crawford and Sobel, 1982](#); [Gilligan and Krehbiel, 1989](#); [Austen-Smith, 1990, 1993](#)). Work on verifiable disclosure asks instead when the unraveling logic of full revelation breaks down ([Milgrom, 1981](#); [Dziuda, 2011](#)). More recent contributions locate frictions inside the advisory relationship itself: [Denisenko, Hafer and Landa \(2026\)](#) show that a more competent advisor anticipates a larger policy movement away from her own ideal point and therefore reveals less, so that a leader can prefer an advisor of limited competence. What these models share is a setting with a sender and a receiver and no audience. The difficulty they identify lies with the advisor — in her incentives to distort, withhold, or decline to acquire information — and the decision-maker’s problem is to extract what she knows.

The present model shuts down every one of these frictions. The advisor is commonly known to share the voter’s preferences, so truthful reporting is not an assumption papering over a transmission problem, but the aligned-preferences case in which that problem does not arise; her signal is exogenous; and the incumbent retains authority over the policy. What is added is a third party. The voter observes that advice was sought and what was done, but not what was said, and her inference about the incumbent’s type determines his tenure. The accountability failure identified below therefore cannot be attributed to any defect in the advisory relationship, because in this model there is none. It is a failure that appears only once advice is embedded in an electoral setting.

That setting comes from the literature on political agency, in which voters use an incumbent’s actions in office to infer whether he shares their interests and allocate support accordingly ([Ferejohn, 1986](#); [Maskin and Tirole, 2004](#); [Ashworth, 2012](#)). Accountability is thus an inferential problem: an action disciplines only insofar as it separates congruent from non-congruent incumbents, and an institution matters through what it allows voters to conclude from the same observable choice. A parallel literature on delegation studies the agency relationship between politicians and expert bureaucrats, trading expertise against the risk of bureaucratic drift ([Epstein and O’Halloran, 1999](#); [Huber and Shipan, 2002](#); [Bendor and](#)

Meirowitz, 2004; Gailmard and Patty, 2012). Fox and Jordan (2011) bring these together and provide the closest formal antecedent to this paper, showing that a politician who delegates to a bureaucrat of uncertain type might gain plausible deniability for policies that harm voters. The mechanism here differs in what it requires. The advisor’s alignment is common knowledge, her report is truthful, and the incumbent decides throughout; the only hidden object is the content of the recommendation. Cover thus needs neither bureaucratic drift, nor captured expertise, nor a transfer of authority — a point that matters because expert consultation, unlike delegation, characteristically leaves the official holding the decision.

The model also speaks to works on pandering and career concerns. That literature shows that reputational incentives can lead officials to take the action the public expects rather than the one they believe correct, and that the distortion is driven by the official’s desire to appear congruent rather than by any defect in his information (Canes-Wrone, Herron and Shotts, 2001*a*; Maskin and Tirole, 2004; Prat, 2005; Fox and van Weelden, 2012). The Type-I equilibria below have this structure, with a twist worth noting. The congruent incumbent does not distort his action to signal type; he declines to acquire, or declines to act on, information that would move him off the action the voter expects. What is sacrificed is not the policy he would otherwise choose but the informational input to that choice, so the distortion is invisible in the record: an outside observer sees an incumbent consulting an advisor and adopting the *ex ante* sensible policy, which is also what a fully instrumental use of advice would produce. The paper’s two classes of equilibrium thus correspond to two familiar pathologies. Where advice is not used, the model reproduces pandering. Where it is used, pandering disappears, and cover takes its place. Neither is available without the other, and which obtains turns on a single condition.

Finally, the paper connects to the literature concerning experts and expertise in public policy. Since the knowledge-utilization scholarship of the late 1970s, scholars have distinguished the instrumental use of expertise from its conceptual, symbolic, and political uses (Weiss, 1979; Knorr, 1977; Amara, Ouimet and Landry, 2004; Blum and Straßheim, 2024). Boswell (2009) shows that organizations often invoke expert knowledge not simply to learn, but to lend authority to preferences they already hold. Related work on blame avoidance emphasizes that expert commissions and advisory processes can help political actors absorb or redirect the costs of unfavorable reforms (Weaver, 1986; Hood, 2011). This literature has identified the symbolic use of expertise as an important empirical phenomenon, but it has generally treated it descriptively. The model below provides microfoundations for the symbolic and political uses of expertise.

In the model, instrumental and symbolic uses are not properties of two different advisory institutions. They are different uses of the same observable act of consultation. In an

equilibrium in which the congruent consults and complies with an unfavorable recommendation, he does so to improve policy. A non-congruent incumbent consults to be seen consulting and then uses the resulting ambiguity as cover. Because the recommendation is private, voters cannot distinguish genuine deference from strategic invocation. The symbolic use of expertise is therefore not simply the act of ignoring available knowledge. It is the appropriation of credibility generated by others' genuine reliance on expertise. In this sense, the symbolic use of advice is parasitic on its instrumental use: cover exists only because some incumbents truly follow the advisor.

This point connects the formal model to broader debates over technocracy and democratic accountability. The promise of expertise is that informed advisors can improve public decisions; the democratic concern is that expert authority may make those decisions harder for citizens to evaluate and control (Fischer, 1990; Weingart, 1999; Caramani, 2017; Bertsou and Caramani, 2022). The model identifies a precise mechanism behind this tension. The accountability problem does not arise because experts are disloyal or incompetent. It arises even when the advisory relationship is in perfect working order. Expertise can erode accountability precisely because it is credible: when a well-intentioned incumbent sometimes follows expert advice into an unfavorable policy, a non-congruent incumbent can use consultation to make the same policy appear less revealing of bias.

3 The Model

3.1 The set-up

I consider a model involving an incumbent politician (he) and a representative voter (she). The incumbent politician must choose one of two policy actions, $a \in \{0, 1\}$, on behalf of the representative voter. One can interpret $a = 0$ as “retain the status quo” or “do nothing” and $a = 1$ as “implement a reform” or “impose a lockdown”. The voter’s welfare depends not on the policy directly but on the outcome it generates, which is jointly determined by the policy and an underlying state of the world $\omega \in \{0, 1\}$. If the incumbent’s policy matches the state, the result is favorable to the voters. However, if the incumbent fails to match the state, the result is detrimental to the voter. The following utility function summarizes these payoffs

$$u_V(a, \omega) = \begin{cases} 1 & \text{if } a = \omega, \\ 0 & \text{otherwise.} \end{cases} \quad (1)$$

The voter's payoff is $u_V(a, \omega)$, so her preferences over policies are state-contingent. No player observes ω directly. The common prior probability that $\omega = 1$ is $p \in (\frac{1}{2}, 1)$, so $a = 1$ is the ex ante most favorable policy; the policy most likely to benefit the voter.

The politician is one of two types. A *congruent* politician shares the voter's state-contingent policy preferences. His total payoff is

$$u_V(a, \omega) + \nu \cdot \mathbb{1}\{\text{reelected}\}, \quad (2)$$

where $\nu > 0$ is the value he attaches to holding office. A *non-congruent* (*non-congruent*) politician has state-independent preferences non-congruent toward $a = 0$: he receives private rents of r whenever $a = 0$ is implemented and 0 otherwise. His total payoff is

$$r \cdot \mathbb{1}\{a = 0\} + \nu \cdot \mathbb{1}\{\text{reelected}\}. \quad (3)$$

The rents r are independently drawn from a distribution $G(r)$ on the interval $[r, \bar{r}]$. The politician's type is his private information. The voter's prior probability that the politician is congruent is $\pi \in (0, 1)$.

The game begins with the incumbent politician deciding whether to consult ($c = 1$) an expert advisor before acting or to proceed without seeking advice ($c = 0$). If consulted, the advisor privately observes an informative signal $s \in \{0, 1\}$ about the state ω with $\Pr(s = \omega \mid \omega) = q \in (p, 1)$ and truthfully transmits the recommendation $m = s$ to the politician. While the incumbent politician observes m , the voter does not. The voter only observes whether the politician consulted and the policy he has implemented. The restriction $q > p$ ensures that the voter would prefer the politician to follow the advisor's recommendation rather than act on the prior alone. Two features of the advisor deserve emphasis. First, the advisor is commonly known to share the voter's preferences and to report truthfully: the friction is not who the advisor is or whether she misreports, but only that the content of her recommendation is unobserved. Second, the incumbent retains full authority over the policy; the advisor informs the choice but does not make it².

Once the incumbent has implemented a policy (whether or not he has consulted), an election is held in which the politician either retains or loses office. I assume that when the voter's posterior belief assigns probability μ to the politician being congruent, he is reelected with probability $F(\mu)$, where $F : [0, 1] \rightarrow [0, 1]$ is a smooth, strictly increasing function³.

²This distinguishes this model from delegation models

³Following [Cameron and Gordon \(2023\)](#), I model the reelection probability as $F(\mu - \alpha_c)$, where μ is the voter's posterior about the incumbent and α_c is her prior belief about challenger quality. Since α_c is a fixed parameter, I write $F(\mu)$. This formulation can be microfounded by a second period of policymaking in which the voter retains the incumbent if and only if her posterior μ exceeds a randomly drawn challenger quality

3.2 Timing and information

1. Nature draws the incumbent's type, the rent r , when non-congruent, and the state ω . The incumbent observes his type; no player observes ω .
2. The incumbent decides whether to consult ($c = 1$) or not ($c = 0$). This is observable. If $c = 0$, the game moves to step 4.
3. If $c = 1$, the advisor observes s and sends the private message $m = s$ to the incumbent.
4. The incumbent chooses $a \in \{0, 1\}$.
5. The voter observes (c, a) , forms a posterior $\mu \in [0, 1]$ about the incumbent's type, and the incumbent is reelected with probability $F(\mu)$.
6. The outcome materializes and payoffs are realized.

I measure accountability on two dimensions. The first is an *extensive* margin (deterrence): the proportion of non-congruent incumbents who are deterred from implementing the unfavorable policy. The second is an *intensive* margin (punishment): the probability of reelection after implementing the ex-ante unfavorable policy.

3.3 Notation

Symbol	Interpretation
$\omega \in \{0, 1\}$	state of the world
$a \in \{0, 1\}$	policy action
$p \in (\frac{1}{2}, 1)$	prior probability that $\omega = 1$
$s \in \{0, 1\}$	advisor's signal of the state
$q \in (p, 1)$	accuracy of the advisor's signal
$r \in [\underline{r}, \bar{r}]$	non-congruent rent from $a = 0$
$\nu > 0$	value of office
$\pi \in (0, 1)$	prior probability that the politician is congruent
$\mu \in [0, 1]$	voter's posterior that the politician is congruent
$F(\mu)$	reelection probability given posterior μ

Table 1: Notation

$\rho \sim F$, giving $F(\mu)$. See [Maskin and Tirole \(2004\)](#) and [Fox and Jordan \(2011\)](#) for a formal treatment.

3.4 Equilibrium concept and refinement

The solution concept is the perfect Bayesian equilibrium. I restrict attention to equilibria in which off-path beliefs satisfy a refinement in the spirit of the D1 criterion of [Banks and Sobel \(1987\)](#). Some discipline over beliefs is needed because consultation and policy are both observable while the recommendation is private. Without it, the game supports a large set of perfect Bayesian equilibria, since arbitrary hostile beliefs after unexpected histories can deter almost anything. One could shut down all consultation by letting the voter read any visit to the advisor as a mark of bias, or force universal consultation by reading any direct policy the same way. Neither construction respects the incentives that actually separate the types at those histories.

Because the model has a continuum of non-congruent rent types, it is convenient to state the refinement through critical reputations. Fix a candidate equilibrium and an off-path public history $h = (c, a)$. A *sure* deviation to h generates h with probability one. Let U_t denote the equilibrium payoff of type t (with $t = C$ for the congruent type and $t = r$ for a non-congruent type with rent r), and let $u_t(h)$ denote the policy component of t 's payoff from a sure deviation to h , net of reelection. The critical reputation $x_t(h)$ solves

$$U_t = u_t(h) + \nu x_t(h),$$

and is the reelection probability that would leave t indifferent between its equilibrium payoff and the deviation. A type with lower $x_t(h)$ is willing to produce h for a wider set of voter responses, so D1 attributes h to the type or types with the smallest critical reputation.

Two comments. First, the restriction is applied *separately* to each public history that can be produced with certainty. A deviation strategy that generates several publicly distinguishable histories is not treated as a single off-path signal; once a belief has been fixed at each realized history, sequential rationality is checked against every deviation, sure or stochastic. Second, D1 pins the support of an off-path belief but generally not the relative weights of the types that tie. When the minimizing set contains the congruent type and a subset of non-congruent types, I assign posterior probabilities by restricting the equilibrium distribution of types to that minimizing set and renormalizing. Thus, each surviving type retains its relative equilibrium weight, while all types outside the minimizing set receive zero probability. I call this the *tie-preserving* belief;

Definition 1 (D1-refined equilibrium). A perfect Bayesian equilibrium is D1-refined if, after every off-path public history $h = (c, a)$, the voter's posterior assigns positive probability only to the type or types whose critical reputation after h is minimal, and, whenever the congruent type and a subset of non-congruent types both minimize, retains the tie-preserving belief.

4 Analysis

To focus on interesting cases, I make the following assumption.

Assumption 1. The rent support satisfies

$$\underline{r} < \nu[F(1) - F(\pi)] \quad \text{and} \quad \bar{r} > \nu[F(\pi) - F(0)].$$

The two inequalities keep both kinds of incumbent behavior on the table. The first says that even under the sharpest reputational separation the model can generate, some non-congruent incumbents would rather give up the rent than take the electoral hit; the second says that some incumbents value the rent enough to accept that hit. Together, they rule out the uninteresting corners in which every non-congruent incumbent mimics or none does, and they guarantee that every cutoff appearing below is interior.

4.1 The no-advice benchmark

I first analyze the benchmark in which advice is not available, so the incumbent chooses $a \in \{0, 1\}$ directly, and the voter observes only the policy. Because she conditions on the action alone, the relevant posteriors are

$$\mu_1 \equiv \Pr(t = C \mid a = 1), \quad \mu_0 \equiv \Pr(t = C \mid a = 0),$$

the probabilities that the incumbent is congruent after each action.

Since $p > 1/2$, the voter's expected payoff is p from $a = 1$ and $1 - p$ from $a = 0$: absent further information, she wants $a = 1$, and the action a non-congruent incumbent prefers is the one that is ex ante least likely to serve her.

I show in the appendix that the benchmark model has a unique D1-refined equilibrium in which the congruent politician always implements $a = 1$, and non-congruent politicians play a cut-off strategy.

Proposition 1. *Under Assumption 1, the no-advice game has a unique D1-refined equilibrium. The congruent incumbent chooses $a = 1$; there is a unique threshold $r^* \in (\underline{r}, \bar{r})$, solving*

$$r^* = \nu[F(\mu_1) - F(0)], \tag{4}$$

where

$$\mu_1 = \mu_1^B(r^*) = \frac{\pi}{\pi + (1 - \pi)G(r^*)}, \tag{5}$$

such that a non-congruent incumbent chooses $a = 1$ if $r < r^$ and $a = 0$ if $r > r^*$; and a politician who chooses $a = 0$ is reelected with probability $F(0)$.*

To see the logic, notice that in any equilibrium in which the congruent incumbent chooses $a = 1$, he does so because $a = 1$ is both the action his own preferences favor under the prior and the action that maximizes his standing with the voter; he has nothing to trade off. The non-congruent incumbent does. Choosing $a = 0$ secures the rent but reveals bias: since no congruent incumbent ever chooses it, the voter who observes $a = 0$ learns his type with certainty and reelects him with the lowest probability, $F(0)$. Choosing $a = 1$ sacrifices the rent but buys him the reputation of an incumbent who might well be congruent. The threshold r^* is simply the marginal type; the rent at which those two options break even. The cutoff structure follows: incumbents with more at stake in $a = 0$ than r^* take the rent and the punishment, and the rest mimic.

The benchmark therefore delivers accountability in its strongest available form, along the two margins the paper tracks. On the extensive margin, a mass $G(r^*)$ of non-congruent incumbents is deterred outright: they give up the rent and act as the voter would have them act. On the intensive margin, every incumbent who is not deterred is fully identified as non-congruent and punished maximally: that is, he is reelected with probability $F(0)$. What makes maximum punishment possible is that the voter's inference problem is trivial here: $a = 0$ admits exactly one explanation, so it is completely attributed to its author.

In the benchmark, then, the policy that is less likely to serve the voter's interest under the prior also carries the greatest electoral cost. It is against this benchmark — full deterrence up to r^* , full punishment below it — that the effect of private advice is measured.

4.2 When advice is available

I now consider the whole game in which the politician may consult the advisor before acting. Consultation and the eventual policy are both publicly observed, so the incumbent produces one of four histories, $(c, a) \in \{0, 1\}^2$. Let μ_{ca} the voter's posterior after the politician chooses $c \in \{0, 1\}$ and $a \in \{0, 1\}$.

Before analyzing how the politician behaves in equilibrium, it is important to clarify how the voter would like him to behave. Consulting the advisor and complying with his recommendations delivers the voter an expected payoff of q ; choosing $a = 1$ directly gives p ; choosing $a = 0$ directly delivers $1 - p$. An incumbent who consults but then ignores what he is told delivers exactly p or $1 - p$ (the same as if he had never consulted). What the voter values, then, is not the act of consulting but compliance with what consultation reveals: she wants the recommendation obtained and obeyed, and consultation that goes unheeded is

worth nothing to her. It turns out that what separates one equilibrium outcome from another is whether the congruent politician consults and follows the recommendation $m = 0$.

Let σ be the probability that the congruent incumbent consults the advisor, and let

$$\tau \equiv \Pr(a = 0 \mid c = 1, m = 0, C)$$

denote the probability that he complies when the recommendation is $m = 0$. Compliance after $m = 1$ is automatic — a congruent incumbent told $m = 1$ sets $a = 1$ with probability one⁴. What matters for everything that follows is the product

$$\kappa \equiv \sigma\tau,$$

the *effective compliance rate*: the probability that a recommendation for $a = 0$, should the advisor's signal produce one, is both obtained and honored. Letting $\beta = p(1 - q) + (1 - p)q$ denote the unconditional probability that $m = 0$, the probability that a congruent incumbent actually arrives at the history $(1, 0)$ is

$$\lambda \equiv \beta\kappa.$$

The characterization below is stated in terms of λ , but κ is the substantive object: it is the extent to which the advisor's expertise is genuinely brought to bear on policy, and, as will become clear, it is also the single quantity that governs how far accountability erodes.

Proposition 2. *Under Assumption 1, there is a unique $\tilde{r}(\lambda)$, continuous and strictly decreasing in λ with $\tilde{r}(0) = r^*$, such that every D1-refined equilibrium of the advice game is of one of two types.*

Type I (non-compliance $\lambda = 0$).

- (a) *The congruent incumbent chooses $a = 1$ with probability one, whether or not he consults.*
- (b) *Non-congruent incumbents with $r < r^*$ choose $a = 1$; those with $r \geq r^*$ choose $a = 0$, and are indifferent between doing so directly and doing so after consulting.*

Type II (compliance, $\lambda > 0$).

- (a) *The congruent incumbent consults with probability $\sigma > 0$; he chooses $a = 1$ when he does not consult and when $m = 1$, and chooses $a = 0$ with probability $\tau > 0$ when $m = 0$.*

⁴I show in the appendix that a congruent incumbent who receives $m = 1$ chooses $a = 1$ with probability one.

(b) *Non-congruent incumbents with $r < \tilde{r}(\lambda)$ choose $a = 1$; those with $r \geq \tilde{r}(\lambda)$ consult and choose $a = 0$ with certainty.*

In Type-I equilibria, henceforth *non-compliance* equilibria, the advisor's expertise is never brought to bear: whatever the incumbent may or may not have been told, $a = 1$ is what he does. Two quite different behaviors produce this ($\lambda = 0$). Under *avoidance* ($\sigma = 0$) the congruent incumbent does not consult at all and always implements $a = 1$ directly. Under *pandering* ($\sigma > 0, \tau = 0$), he sometimes asks, but sets aside any recommendation that would move him off $a = 1$, preferring the safety of the prior to the reputational risk of appearing to have been talked into something unfavorable (Canes-Wrone, Herron and Shotts, 2001b; Maskin and Tirole, 2004). To an observer, these look nothing alike — one leaves a visible record of consultation and the other does not — but they are strategically identical, because in neither case does consultation carry any information about what was subsequently done. The voter learns that asking does not change the policy.

The consequence is that non-compliance equilibria reproduce the benchmark exactly. Since the congruent incumbent never chooses $a = 0$, that action is again explained only by non-congruence, and it draws the posterior zero whether or not the advisor was consulted first. The same mass $G(r^*)$ of non-congruent incumbents is deterred, and any incumbent who chooses $a = 0$ — with a consultation on the record or without one — is reelected with probability $F(0)$. Consultation that changes nothing insulates no one.

Matters are different in Type-II equilibria, henceforth *compliance* equilibria. Here, the congruent incumbent sometimes lets the advisor overturn his own default: recommended $m = 0$, he chooses $a = 0$ with probability $\tau > 0$. This is exactly the behavior that gives the advisor her value: it raises the congruent incumbent's expected accuracy from p the accuracy of the prior, to $p + \kappa(q - p)$. Compliance is the only source of that improvement, since neither consulting without acting on the recommendation nor acting without consulting makes the action dependent on the state ω . It is also the behavior that changes what the voter can conclude from what she sees.

To understand how the non-congruent incumbent escapes maximal punishment in a compliance equilibrium, first notice that although $a = 0$ is ex ante unlikely to serve the voter, there is always a positive probability that the advisor's signal points toward it. A congruent incumbent who consults and complies will therefore sometimes choose $a = 0$ for entirely honest reasons. Since the voter observes that the advisor was consulted and that $a = 0$ was chosen, but never what was recommended, the history $(1, 0)$ is generated by two very different incumbents: a congruent one who was told $m = 0$ and obeyed, and a non-congruent one who consulted purely for appearances and took his rent regardless of what he heard. The two are observationally identical — the same visit to the advisor, the same action — the voter

cannot tell which incumbent she is looking at. The voter must, therefore, assign positive probability to congruence after $(1, 0)$. That is, $\mu_{10} > 0$, and the incumbent who shirks along this route is reelected with probability strictly above $F(0)$.

One can think of the non-congruent incumbent in a compliance equilibrium as enjoying what [Fox and Jordan \(2011\)](#) call "plausible deniability". The voter knows that incumbents who share her interests sometimes seek advice and act on it. She knows that $a = 0$ is unlikely to be what her interests require, but also that an accurate advisor will occasionally conclude that it is. So when she observes consultation followed by $a = 0$, she cannot rule out that she is looking at an incumbent who did precisely what she would have wanted. The non-congruent incumbent needs to say nothing more than this: that he consulted, and that he did what he was told. The claim is false and unfalsifiable at once, and it is credible for the same reason the appeal to having "followed the science" is credible — because some incumbents really do follow it.

Note that the deniability attaches to the path of play, not only to the action. By [Proposition 2](#), the congruent incumbent never chooses $a = 0$ without consulting; thus, an incumbent who chooses $a = 0$ directly still identifies himself as non-congruent and is still reelected with probability $F(0)$, exactly as in the benchmark. Implementing $a = 0$ after consultation is ambiguous. This is why every non-congruent incumbent who shirks in a compliance equilibrium shirks after consulting, and why he does so with certainty. The consultation makes no difference to what he learns — he has no use for the recommendation, since he will choose $a = 0$ whatever it says — and no difference to the voter's payoff, since the action and the outcome are the same either way. In that sense, it is purely cosmetic. But it is free, and it converts a fully attributed act into an ambiguous one. This deniability, in turn, changes the calculation of every non-congruent incumbent, not only those who were going to shirk anyway. Consultation followed by $a = 0$ now provides *cover*, and the availability of cover disturbs the extensive margin. We have the following result.

Corollary 1 (Accountability). Let κ be the effective compliance rate. If $\kappa = 0$, both accountability margins coincide with the benchmark: a mass $G(r^*)$ of non-congruent incumbents is deterred, and any incumbent who chooses $a = 0$ is reelected with probability $F(0)$. If $\kappa > 0$, then

$$\underbrace{G(r^*) - G(\tilde{r}(\beta\kappa))}_{\text{deterrence}} > 0, \quad \text{and} \quad \underbrace{F(\mu_{10}(\tilde{r}(\beta\kappa)))}_{\text{punishment}} > F(0),$$

and both are strictly increasing in κ .

Since $\tilde{r}(\lambda) < r^*$, whenever $\lambda > 0$ a mass $G(r^*) - G(\tilde{r}(\beta\kappa))$ that would have set $a = 1$ in the benchmark now shirks under cover. An incumbent who consults and sets $a = 0$ is reelected with probability $F(\mu_{10}(\beta\kappa)) > F(0)$. Not only do more non-congruent incumbents

choose $a = 0$ than would have done so without an advisor; those who do are punished less than they would have been.

Both movements share a common source, and it is κ . Hold the cutoff fixed and raise the effective compliance rate. The congruent incumbent then appears less often at $(1, 1)$ and more often at $(1, 0)$, so μ_{11} falls while μ_{10} rises. Each movement narrows the reputational gap between mimicking and shirking-under-cover, which is precisely the gap the marginal non-congruent incumbent is weighing against his rent. A narrower gap means a smaller rent suffices to make shirking worthwhile, and the cutoff falls. Both margins are thus increasing in κ : the more genuinely the advisor's expertise is used, the further discipline erodes, and the erosion vanishes only in the limit $\kappa = 0$, which is non-compliance.

Two features of the cover deserve comment. First, it is never complete. In any compliance equilibrium $\mu_{10} < \mu_{11}$: non-congruent incumbents are over-represented at $(1, 0)$ relative to $(1, 1)$, since a congruent incumbent reaches $(1, 0)$ only through the $\beta\kappa$ fraction of cases in which an unfavorable signal both arises and is obeyed, whereas a shirking incumbent reaches it always. Consultation therefore softens the punishment for shirking but never erases it, and the cutoff can never be pushed all the way to $\bar{r}$ — were it to be, the pool at $(1, 0)$ would become entirely congruent, and the cover would cease to be worth seeking, which is self-defeating.

The second is that the cover is parasitic. A non-congruent incumbent can appropriate the credibility of expert advice only because congruent incumbents sometimes genuinely defer to it. If κ were zero — if the advisor were ignored, or too inaccurate to be worth heeding — then $(1, 0)$ would be no more explicable than $a = 0$ chosen directly, and consultation would purchase nothing. The symbolic use of expertise in this model is therefore not the mere act of invoking experts while ignoring them. It is the appropriation of credibility that others' genuine reliance on expertise has generated. What makes the advisor useful is what makes her usable.

It is worth stating how the mechanism here differs from formal models of delegation ([Fox and Jordan, 2011](#)). In Fox and Jordan (2011), the politician hands the decision to a bureaucrat whose type he observes and the voter does not; plausible deniability arises because the voter cannot tell a delegation meant to serve her from one meant to circumvent her, and part of the escape from blame runs through another actor's authority over the choice. Here, none of those ingredients are present. The advisor's alignment with the voter is common knowledge, she reports her signal truthfully, and the incumbent retains the decision throughout. The voter never observes what was recommended. That alone is enough. The accountability problem is thus not a problem of captured or drifting expertise, nor of transferred authority, nor of consultation as such; it is a problem of advice that is private and genuinely used.

4.3 When does a compliance equilibrium exist?

I now identify the conditions under which a compliance equilibrium exists. By doing so, I characterize the conditions under which an expert advisor provides cover for non-congruent politicians. The existence of compliance equilibria hinges upon the congruent politician's obedience with an unfavorable recommendation. The key comparison is between the policy value of complying with an unfavorable recommendation and its reputational cost.

Conditional on $m = 0$, switching from $a = 1$ to $a = 0$ raises the congruent incumbent's probability of matching the state by

$$\Pr(\omega = 0 \mid m = 0) - \Pr(\omega = 1 \mid m = 0) = \frac{q - p}{\beta}, \quad (6)$$

which I call the *value of compliance*. It increases with the advisor's accuracy and vanishes as $q \rightarrow p$. The cost of complying with $m = 0$ is reputation. Following $m = 0$ moves the incumbent from a history the voter reads well $(1, 1)$ to one she reads less well $(1, 0)$. At effective compliance κ , moving from $a = 1$ to the $a = 0$ after consultation costs

$$\nu \left[F(\mu_{11}(\tilde{r}(\beta\kappa))) - F(\mu_{10}(\tilde{r}(\beta\kappa))) \right] = \tilde{r}(\beta\kappa),$$

where the equality is the marginal non-congruent incumbent's indifference condition. The reputational distance between the two histories is exactly what makes him willing to trade his rent, so the congruent incumbent's compliance cost is the discipline cutoff itself. As the next proposition establishes, a compliance equilibrium exists if and only if the value of compliance exceeds its reputational cost.

Proposition 3. *Under Assumption 1, a compliance equilibrium with effective compliance :*

(a) $\kappa = 1$ exists if and only if

$$\frac{q - p}{\beta} > \tilde{r}(\beta) \equiv \hat{r} \quad (\text{full compliance}).$$

(b) $0 < \kappa < 1$ exists if and only if

$$\frac{q - p}{\beta} = \tilde{r}(\beta\kappa) \quad (\text{partial compliance}).$$

The condition is a statement about the congruent incumbent's willingness to bear an electoral cost in order to get the policy right. Complying raises his chance of matching the state by $\frac{q - p}{\beta}$. It costs him $\tilde{r}(\beta\kappa)$ in reputation. He complies only when the first exceeds

the second. What makes the comparison unusual is that the price is not a free parameter. It is the discipline cutoff itself — the same quantity that measures how far accountability has eroded on the extensive margin. Advice is used only where it is worth paying for, and it is paid for in discipline.

Two implications follow. First, the accountability problem requires an advisor who is informative enough to be worth heeding. As q falls toward p , the value of compliance falls toward zero, and below $\tilde{r}(\beta)$ no compliance is incentive-compatible at any κ ; the congruent incumbent is unwilling to consult and obey, and only non-compliance behavior survives. Concerns about advice undermining accountability are therefore misplaced where the advisory body has little genuine informational advantage: there, consultation may be frequent and visible, but it will not be followed, and what is not followed provides nothing to borrow.

Second, the problem also requires that the incumbent care about policy relative to office. The price of compliance increases with the value of office: $\tilde{r}$ is proportional to ν at given posteriors, so a more ambitious incumbent faces a steeper reputational cost for the same act of obedience. For sufficiently large ν , the condition fails and only non-compliance survives. This inverts the usual worry. The accountability failure is not caused by incumbents who care too much about reelection; those incumbents pander, and pandering leaves both margins intact. It requires incumbents who care enough about getting policy right to accept an electoral cost for doing so, and it is exactly their willingness to accept that cost which supplies cover to incumbents who will accept nothing.

The effect of advisor accuracy on existence is not straightforward because raising q raises both sides at once: better advice is worth more to the congruent incumbent, but it also improves the pool at $(1, 1)$ and thins the pool at $(1, 0)$, which widens the reputational gap he must pay. The two do not move at comparable rates. Since $p > 1/2$, an increase in q makes an unfavorable recommendation rarer, so the price $\tilde{r}(\beta)$ responds weakly, while v rises without bound relative to it. Better advice, therefore, makes a compliance equilibrium easier to sustain, and the same increase that sustains it also raises the cutoff $\tilde{r}(\beta\kappa)$ within it. Accuracy has opposing effects on the two questions the paper asks: it makes cover more likely to exist, and less damaging when it does.

5 Voter welfare

The accountability result raises a further question. If private advice can weaken electoral discipline, should voters prefer to restrict access to such advice or require recommendations to be disclosed? The worry echoes a familiar concern about institutions that let officials

escape blame for the choices they make (Fox and Jordan, 2011): if non-congruent incumbents can hide behind consultation, restricting consultation might seem to serve voters better. In the present model, that intuition is only half right. Private advice weakens accountability precisely when it is used — when $\kappa > 0$ but the same compliance that supplies cover for non-congruent incumbents also improves the choices of a congruent one, and which effect dominates is not settled by the accountability result on its own.

The relevant measure of voter welfare is her expected policy payoff. The benchmark is the no-advice equilibrium. I compare this benchmark to an equilibrium of the advice game with an effective compliance rate κ and cutoff $\tilde{r}(\beta\kappa)$. This comparison separates the two effects of advice. On the one hand, advice creates an expertise gain: congruent incumbents sometimes replace prior-based policy with a more accurate recommendation. On the other hand, advice creates a discipline loss: some non-congruent incumbents who would have mimicked the congruent type in the benchmark now choose the ex ante unfavorable policy under advisory cover.

In a compliance equilibrium with compliance rate $\kappa > 0$ and cutoff $\tilde{r}(\beta\kappa)$, the congruent incumbent's expected policy accuracy is $p + \kappa(q - p)$: he acts on the prior except when an unfavorable recommendation arises and is obtained and obeyed, an event with probability $\beta \cdot \kappa$, each occurrence of which raises his accuracy by the value of compliance $(q - p)/\beta$. Non-congruent incumbents below the cutoff choose $a = 1$, correct with probability p ; those above choose $a = 0$, correct with probability $1 - p$.

Proposition 4. *Let $W(\kappa)$ be voter welfare in an advice-game equilibrium with effective compliance κ and cutoff $\tilde{r}(\beta\kappa)$, and let W_B be welfare in the no-advice benchmark. Then*

$$W(\kappa) - W_B = \underbrace{\pi\kappa(q - p)}_{\text{expertise gain}} - \underbrace{(1 - \pi)[G(r^*) - G(\tilde{r}(\beta\kappa))]}_{\text{discipline loss}}(2p - 1). \quad (7)$$

The first term in (7) is the improvement in policy accuracy from a congruent incumbent. The second is the loss from the non-congruent types between $\tilde{r}(\beta\kappa)$ and r^* who switch from $a = 1$ to $a = 0$; each switch reduces expected accuracy from p to $1 - p$, a loss of $2p - 1$. What the expression also makes clear is that the two effects cannot be separated. Both are increasing in κ : it is the same act of compliance that improves the congruent incumbent's decision and makes (1,0) an action a congruent incumbent might plausibly have taken. There is no configuration of the model in which the voter obtains the expertise gain without incurring the discipline loss, because they are the same behavior described from her two different points of view. Accountability and welfare should therefore not be conflated. Every equilibrium with $\kappa > 0$ weakens accountability; some of those equilibria nonetheless leave the voter better off

than she would be with no advisor at all, and some do not.

The accuracy of the advisor moves both terms in the same direction. A higher q raises the expertise gain directly, through $(q - p)$; and it raises the advisory cutoff, which shrinks the displaced mass $G(r^*) - G(\tilde{r}(\beta\kappa))$ while leaving the benchmark cutoff r^* untouched⁵. The welfare gap $W(\kappa) - W_B$ is therefore increasing in q on both counts. This cuts against the natural reading of the accountability result. One might expect a sharper advisor to supply better cover and so to do more damage. The opposite holds: a sharper advisor makes the congruent incumbent's arrival at $(1, 0)$ rarer relative to the non-congruent incumbent's, so the covered pool deteriorates, and the reputational distance between mimicking and shirking widens. Cover is most corrosive when it is cheapest to obtain, which is when the advisor is barely worth consulting.

The voter's prior about incumbent congruence pulls in no clear direction. A higher π raises the benchmark cutoff unambiguously, since mimicking now buys a better reputation. In a compliance equilibrium it improves both posteriors at once: $(1, 1)$ looks better, but so does $(1, 0)$, so mimicking and covered shirking become more attractive together. Which dominates depends on the curvature of F , the density G near the cutoff, and where the cutoff sits. The same indeterminacy carries into welfare, where π additionally scales the two terms of Proposition 4 in opposite directions.

A caveat applies to all of these comparisons. Over the intermediate range of accuracy identified in Proposition 2, a non-compliance equilibrium ($\kappa = 0$), a partial compliance equilibrium ($0 < \kappa < 1$), and the full-compliance equilibrium ($\kappa = 1$) coexist, so a change in primitives has no well-defined effect on outcomes unless one fixes which equilibrium is played. The statements above hold within a fixed selection: the benchmark against the full-compliance equilibrium, with parameter changes restricted to those preserving its existence. Along the partial-compliance family, the comparison is different and simpler, since the cutoff equals the value of compliance, $\tilde{r} = (q - p)/\beta$, which is strictly increasing in q ; better advice raises the cutoff directly rather than through the posteriors.

6 Discussion

The role of ‘the expert’ in public decision-making has been a cornerstone of modern policy-making. While scholars acknowledge the substantive benefits of expertise, critics fear that the invocation of experts can undermine accountability. However, these critics are relatively silent on how, and under what conditions, experts in politics undermine accountability. This paper addresses this limitation by providing a theoretical framework tailored to identifying

⁵The benchmark has no advisor, so q does not enter it.

exactly when these concerns are warranted.

The analysis suggests that two conditions are necessary for experts to undermine accountability. First, the expert's recommendation must be undisclosed, so that the voter sees an incumbent consulting and acting but cannot tell whether the action honored the advice or defied it. Second, the advice must be informative enough that a congruent incumbent is willing to bear a reputational cost to follow it, which occurs when the accuracy gained from compliance exceeds the discipline the equilibrium prices it at. The first condition makes genuine and strategic consultation observationally equivalent; the second ensures that genuine consultation occurs often enough for the equivalence to matter. Together, they mean that cover is borrowed rather than manufactured. A biased incumbent's claim to have followed the evidence is credible only because other incumbents follow it in earnest, and the symbolic use of expertise is therefore parasitic on its instrumental use.

This locates the risk somewhere other than where public debate usually places it. The mechanism does not require a captured advisor, a dishonest one, or a transfer of authority away from the elected official; it requires only that the advisory relationship function well and privately. Ceremonial consultation of bodies with little genuine informational advantage is harmless here. What generates cover is a capable advisor whose recommendations are actually followed. Expertise erodes accountability not despite being credible but because it is.

The welfare implications are correspondingly not straightforward. Advice weakens electoral discipline only when some incumbents use it to improve their decisions, so the gain and the loss have a common source and scale together with the rate at which advice is genuinely obeyed. Whether the voter is better off therefore turns on whether the improvement in congruent incumbents' policies outweighs the shirking that their compliance makes possible—and the two can come apart in either direction. An arrangement that supplies real information can leave voters worse off, and an arrangement that measurably weakens accountability can leave them better off. Accountability and welfare are distinct standards here, and an institution can fail one while satisfying the other.

Because the problem stems from private advice, the natural remedy is to disclose the recommendation. While disclosure does not eliminate the pooling that generates cover, it meaningfully relocates it. An incumbent advised to take one action but chooses another reveals his true type, since no congruent incumbent would do this. What remains pooled are cases where the incumbent follows the advisor's recommendation—in these situations, a biased incumbent escapes punishment not by manufacturing ambiguity, but simply because circumstances favored him. Thus, disclosure leaves a selection problem: voters cannot always identify the incumbent's type, but it does resolve the control problem. Now, a biased incumbent can only avoid electoral costs when his preferred action coincides with what an

informed advisor would have recommended. The real issue was never the choice itself, but rather when the choice was made against the evidence—and the difficulty of distinguishing that from choices genuinely based on it.

However, this remedy is more limited than it may seem. Disclosure disciplines because a published recommendation certifies what the evidence supports, but this presumes an advisor whose recommendation is genuinely credible. If the advisor shares the incumbent’s bias, publication provides a public warrant rather than a check, and cover returns in a more durable form—the alibi becomes part of the record instead of an unverifiable claim. When incumbents can choose whom to consult, publication incentivizes selecting advisors whose recommendations are predictable, shifting the problem from how advice is used to whom it is sought. If advisory bodies control the scope of their inquiries, anticipating disclosure may shape what they investigate—a dynamic established in non-electoral settings by [Denisenko, Hafer and Landa \(2026\)](#). These issues fall outside the present model, but each would influence whether disclosure is desirable in practice. Making advice observable disables the mechanism identified here, but its broader effects on the advisory relationship remain an open question for future research.

References

- Amara, Nabil, Mathieu Ouimet and Réjean Landry. 2004. “New Evidence on Instrumental, Conceptual, and Symbolic Utilization of University Research in Government Agencies.” *Science Communication* 26(1):75–106.
- Ashworth, Scott. 2012. “Electoral Accountability: Recent Theoretical and Empirical Work.” *Annual Review of Political Science* 15:183–201.
- Austen-Smith, David. 1990. “Information transmission in debate.” *American Journal of political science* pp. 124–152.
- Austen-Smith, David. 1993. “Information and influence: Lobbying for agendas and votes.” *American Journal of Political Science* pp. 799–833.
- Banks, Jeffrey S. and Joel Sobel. 1987. “Equilibrium Selection in Signaling Games.” *Econometrica* 55(3):647–661.
- Bendor, Jonathan and Adam Meirowitz. 2004. “Spatial Models of Delegation.” *American Political Science Review* 98(2):293–310.
- Bertsou, Eri and Daniele Caramani. 2022. “People Haven’t Had Enough of Experts: Technocratic Attitudes among Citizens in Nine European Democracies.” *American Journal of Political Science* 66(1):5–23.
- Blum, Sonja and Holger Straßheim. 2024. Experts and Expertise in Public Policy. In *Encyclopedia of Public Policy*, ed. Minna van Gerven, Christine Rothmayr Allison and Klaus Schubert. Cham: Springer.
- Boswell, Christina. 2009. *The Political Uses of Expert Knowledge: Immigration Policy and Social Research*. Cambridge: Cambridge University Press.
- Calvert, Randall L. 1985. “The value of biased information: A rational choice model of political advice.” *The Journal of Politics* 47(2):530–555.
- Cameron, Charles M and Sanford C Gordon. 2023. “Fire Alarms and Democratic Accountability.” *Accountability Reconsidered: Voters, Interests, and Information in US Policymaking* pp. 221–41.
- Canes-Wrone, Brandice, Michael C. Herron and Kenneth W. Shotts. 2001a. “Leadership and Pandering: A Theory of Executive Policymaking.” *American Journal of Political Science* 45(3):532–550.

- Canes-Wrone, Brandice, Michael C Herron and Kenneth W Shotts. 2001*b*. “Leadership and pandering: A theory of executive policymaking.” *American Journal of Political Science* pp. 532–550.
- Caramani, Daniele. 2017. “Will vs. Reason: The Populist and Technocratic Forms of Political Representation and Their Critique to Party Government.” *American Political Science Review* 111(1):54–67.
- Crawford, Vincent P. and Joel Sobel. 1982. “Strategic Information Transmission.” *Econometrica* 50(6):1431–1451.
- Denisenko, Anna, Catherine Hafer and Dimitri Landa. 2026. “Competence and advice.” *American Journal of Political Science* 70(1):21–39.
URL: <https://onlinelibrary.wiley.com/doi/abs/10.1111/ajps.12905>
- Dziuda, Wioletta. 2011. “Strategic argumentation.” *Journal of Economic Theory* 146(4):1362–1397.
- Ehlers, Sarah and Stefan Esselborn, eds. 2022. *Evidence in Action between Science and Society: Constructing, Validating, and Contesting Knowledge*. Taylor & Francis.
- Epstein, David and Sharyn O’Halloran. 1999. *Delegating Powers: A Transaction Cost Politics Approach to Policy Making under Separate Powers*. New York: Cambridge University Press.
- Ferejohn, John. 1986. “Incumbent Performance and Electoral Control.” *Public Choice* 50(1–3):5–25.
- Fischer, Frank. 1990. *Technocracy and the Politics of Expertise*. Newbury Park, CA: Sage.
- Fox, Justin and Richard van Weelden. 2012. “Costly Transparency.” *Journal of Public Economics* 96(1–2):142–150.
- Fox, Justin and Stuart V. Jordan. 2011. “Delegation and Accountability.” *Journal of Politics* 73(3):831–844.
- Gailmard, Sean and John W. Patty. 2012. “Formal Models of Bureaucracy.” *Annual Review of Political Science* 15:353–377.
- Gilligan, Thomas W. and Keith Krehbiel. 1987. “Collective Decisionmaking and Standing Committees.” *Journal of Law, Economics, and Organization* 3(2):287–335.

- Gilligan, Thomas W. and Keith Krehbiel. 1989. "Asymmetric Information and Legislative Rules with a Heterogeneous Committee." *American Journal of Political Science* 33(2):459–490.
- Hood, Christopher. 2011. *The Blame Game: Spin, Bureaucracy, and Self-Preservation in Government*. Princeton: Princeton University Press.
- Huber, John D. and Charles R. Shipan. 2002. *Deliberate Discretion? The Institutional Foundations of Bureaucratic Autonomy*. New York: Cambridge University Press.
- Knorr, Karin D. 1977. Policymakers' Use of Social Science Knowledge: Symbolic or Instrumental? In *Using Social Research in Public Policy Making*, ed. Carol H. Weiss. Lexington, MA: Lexington Books pp. 165–182.
- Krishna, Vijay and John Morgan. 2001. "A model of expertise." *The Quarterly Journal of Economics* 116(2):747–775.
- Maskin, Eric and Jean Tirole. 2004. "The Politician and the Judge: Accountability in Government." *American Economic Review* 94(4):1034–1054.
- McKee, Martin, Kara Hanson and Kamran Abbasi. 2022. "Guided by the Science? Questions for the UK's Covid-19 Public Inquiry." *BMJ* 378.
- Milgrom, Paul R. 1981. "Good news and bad news: Representation theorems and applications." *The Bell Journal of Economics* pp. 380–391.
- Prat, Andrea. 2005. "The wrong kind of transparency." *American economic review* 95(3):862–877.
- Weaver, R. Kent. 1986. "The Politics of Blame Avoidance." *Journal of Public Policy* 6(4):371–398.
- Weingart, Peter. 1999. "Scientific Expertise and Political Accountability: Paradoxes of Science in Politics." *Science and Public Policy* 26(3):151–161.
- Weiss, Carol H. 1979. "The Many Meanings of Research Utilization." *Public Administration Review* 39(5):426–431.

Appendix

A The no-advice benchmark

In the benchmark, the incumbent chooses $a \in \{0, 1\}$ directly and the voter conditions on the action alone. If every non-congruent type with rent below r chooses $a = 1$, Bayes' rule gives the posterior after $a = 1$,

$$\mu_1^B(r) \equiv \frac{\pi}{\pi + (1 - \pi)G(r)}, \quad (\text{A1})$$

which is strictly decreasing in r , with $\mu_1^B(\underline{r}) = 1$ and $\mu_1^B(\bar{r}) = \pi$.

Let

$$\Psi^B(r) \equiv \nu[F(\mu_1^B(r)) - F(0)]. \quad (\text{A2})$$

Lemma 1 (Congruent behavior). *In every D1-refined equilibrium of the no-advice game, the congruent incumbent chooses $a = 1$ with probability one, and consequently $\mu_0 = 0$.*

Proof. Suppose first that the congruent type mixes. Indifference requires $p + \nu F(\mu_1) = (1 - p) + \nu F(\mu_0)$, and since $p > \frac{1}{2}$ this forces $F(\mu_0) > F(\mu_1)$, hence $\mu_0 > \mu_1$. But then every non-congruent type strictly prefers $a = 0$, which delivers both the rent and the higher reputation, so only the congruent type chooses $a = 1$ and $\mu_1 = 1 \geq \mu_0$, a contradiction.

Suppose instead that the congruent type chooses $a = 0$ with probability one. Any non-congruent type choosing $a = 1$ would obtain no rent and a posterior of zero, and would strictly prefer to pool at $a = 0$. Thus, all types pool at $a = 0$, the on-path posterior would be π , and $a = 1$ is off-path. A sure deviation to $a = 1$ has critical reputations

$$x_C(1) = F(\pi) - \frac{2p - 1}{\nu}, \quad x_r(1) = F(\pi) + \frac{r}{\nu},$$

so the congruent type is the unique minimizer, and D1 assigns the deviation posterior one. The deviation then yields $p + \nu F(1) > (1 - p) + \nu F(\pi)$, a contradiction.

Hence, the congruent type chooses $a = 1$ with probability one. Any incumbent who chooses $a = 0$ is therefore non-congruent, and $\mu_0 = 0$ whether $a = 0$ is on path (Bayes) or off path (only non-congruent types minimize, since relocating to $a = 0$ preserves the rent). $\square$

Lemma 2 (Existence and uniqueness of the benchmark cutoff). *There is a unique $r^* \in (\underline{r}, \bar{r})$ solving $r = \Psi^B(r)$.*

Proof. Since μ_1^B is strictly decreasing and F strictly increasing, Ψ^B is strictly decreasing, so $h^B(r) \equiv r - \Psi^B(r)$ is strictly increasing and continuous. By (??),

$$h^B(\underline{r}) = \underline{r} - \nu[F(1) - F(0)] < \underline{r} - \nu[F(1) - F(\pi)] < 0,$$

$$h^B(\bar{r}) = \bar{r} - \nu[F(\pi) - F(0)] > 0.$$

The zero is unique and interior. $\square$

Proof of Proposition 1. By Lemma 1, the congruent type chooses $a = 1$ and $\mu_0 = 0$. For a non-congruent type with rent r , the payoff difference between shirking and mimicking is

$$[r + \nu F(0)] - \nu F(\mu_1) = r - \Psi^B(\cdot),$$

which is strictly increasing in r and independent of r through the posterior. Best responses are therefore cutoff rules, and the marginal type solves $r = \Psi^B(r)$, which by Lemma 2 has the unique interior solution r^* . Types $r < r^*$ choose $a = 1$ and types $r > r^*$ choose $a = 0$. Both actions occur with positive probability, so Bayes' rule fixes $\mu_1 = \mu_1^B(r^*)$ and $\mu_0 = 0$, delivering (A1) and the reelection probability $F(0)$ after $a = 0$. Conversely, these strategies and beliefs constitute a D1-refined equilibrium: the congruent type strictly prefers $a = 1$ because it is optimal on policy under the prior and carries the higher posterior, and non-congruent optimality is the cutoff condition. $\square$

B Preliminaries for the advice game

B.1 Signal, compliance, and strategies

A consulted advisor observes s with $\Pr(s = \omega \mid \omega) = q \in (p, 1)$ and transmits $m = s$. Write

$$\beta \equiv \Pr(m = 0) = p(1 - q) + (1 - p)q \in (0, \tfrac{1}{2}). \quad (\text{B1})$$

After $m = 0$, the congruent politician's policy gain from choosing $a = 0$ rather than $a = 1$ is

$$\Pr(\omega = 0 \mid m = 0) - \Pr(\omega = 1 \mid m = 0) = \frac{q - p}{\beta}. \quad (\text{B2})$$

After $m = 1$, the congruent politician's policy gain from choosing $a = 1$ rather than $a = 0$ is

$$\Pr(\omega = 1 \mid m = 1) - \Pr(\omega = 0 \mid m = 1) = \frac{p + q - 1}{1 - \beta} > 0. \quad (\text{B3})$$

Both inequalities follow from $q > p > \frac{1}{2}$.

For later use, write

$$v \equiv \frac{q - p}{\beta}, \quad (\text{B4})$$

For each politician type t , a complete behavioral strategy is

$$s_t = (\sigma_t, \alpha_t), \quad \alpha_t : \{\emptyset, 0, 1\} \rightarrow [0, 1],$$

where $\sigma_t = \Pr(c = 1 \mid t)$; $\alpha_t(\emptyset) = \Pr(a = 1 \mid c = 0, t)$; $\alpha_t(m) = \Pr(a = 1 \mid c = 1, m, t)$, $m \in \{0, 1\}$.

Thus, the strategy specifies consultation, the direct policy, and the policy after every possible recommendation following a consultation. For the congruent type, write $\sigma \equiv \sigma_C$ and $\tau \equiv 1 - \alpha_C(0)$. Lemmas 4 and 9 establish $\alpha_C(\emptyset) = \alpha_C(1) = 1$, so the congruent type's only route to $(1, 0)$ is consultation, an unfavorable recommendation, and positive compliance. The effective compliance rate is $\kappa \equiv \sigma\tau$ and

$$\lambda \equiv \beta\kappa = \beta\sigma\tau \in [0, \beta]. \quad (\text{B5})$$

B.2 Location of types across histories

Lemma 3 (Recommendation independence). *Conditional on consulting, a non-congruent type's optimal action does not depend on the recommendation except under indifference; a signal-contingent rule is never a strict best response.*

Proof. After consultation, a non-congruent type r obtains

$$r\mathbb{1}\{a = 0\} + \nu F(\mu_{1a}).$$

Neither the rent nor the reputational payoff depends on the private recommendation conditional on the public action. The comparison between $a = 0$ and $a = 1$ is therefore identical after $m = 0$ and $m = 1$. If

$$U_S(r) = r + \nu F(\mu_{10}), \quad U_M = \nu F(\mu_{11}),$$

a rule that selects $a = 0$ after recommendations with total probability $\zeta \in [0, 1]$ yields

$$\zeta U_S(r) + (1 - \zeta)U_M \leq \max\{U_S(r), U_M\}.$$

Thus a signal-contingent rule can be optimal only when the two constant rules are tied. $\square$

Lemma 4 (No direct $a = 0$). *In every D1-refined equilibrium, $\alpha_C(\emptyset) = 1$.*

Proof. Suppose the congruent type chose $(c, a) = (0, 0)$ with positive probability. I first show that no history with $a = 1$ could then be on path. Let h_1 be such a history. The congruent

type can generate h_1 by a sure plan with policy payoff p . Optimality of direct $a = 0$ would imply $(1 - p) + \nu F(\mu_{00}) \geq p + \nu F(\mu(h_1))$, and hence $F(\mu_{00}) - F(\mu(h_1)) \geq \frac{2p-1}{\nu} > 0$. Any non-congruent type at h_1 would then strictly prefer direct $a = 0$, which gives both a rent and a higher reelection probability. Therefore, no non-congruent type can be at h_1 . If the congruent type generated h_1 with positive probability, Bayes' rule would give $\mu(h_1) = 1$, which would require $F(\mu_{00}) > F(1)$, a contradiction. Thus, every type chooses $a = 0$ on path.

Now consider the off-path deviation to $(0, 1)$, where the congruent incumbent implements $a = 1$ without advice. The congruent type's critical reelection probability is

$$x_C(0, 1) = F(\mu_{00}) - \frac{2p - 1}{\nu}.$$

A non-congruent type located at either history with $a = 0$ has a critical reelection probability of at least $F(\mu_{00}) + r/\nu$: if both histories with $a = 0$ are on path, optimality for the congruent and non-congruent types forces their posteriors to be equal; if only one is on path, use its posterior. Hence, the congruent type is the unique D1 minimizer and $\mu_{01} = 1$. Direct $a = 1$ then yields $p + \nu F(1) > (1 - p) + \nu F(\mu_{00})$, meaning the congruent incumbent has a profitable deviation; a contradiction. $\square$

Lemma 5 (Location and equalization across $a = 1$ histories). *Fix an equilibrium in which the congruent politician reaches $a = 1$ with positive probability.*

1. *If both $(0, 1)$ and $(1, 1)$ are on path, then*

$$\mu_{01} = \mu_{11}.$$

The masses of non-congruent mimickers at the two histories are proportional to the congruent type's probabilities of reaching them.

2. *If exactly one history with $a = 1$ is on path, every non-congruent mimicker uses that history.*

Consequently, consultation has no independent reputational value across histories with $a = 1$.

Proof. Suppose both $(0, 1)$ and $(1, 1)$ are on path and, contrary to the claim, $\mu_{01} > \mu_{11}$. Because consultation is costless and a non-congruent mimicker receives no rent at either history, every mimicker strictly prefers $(0, 1)$. The history $(1, 1)$ is nevertheless reached by the congruent type, so Bayes' rule would give $\mu_{11} = 1$, contradicting $\mu_{01} > \mu_{11}$. The reverse inequality yields the symmetric contradiction. Hence $\mu_{01} = \mu_{11}$.

For Bayes' rule to yield the same posterior at both histories, the ratio of non-congruent to congruent probability must be the same at each. This is exactly proportional splitting. If only one route with $a = 1$ is used by the congruent type, a positive mass of mimickers at the other route would generate posterior zero there and would strictly relocate to the route containing the congruent type. $\square$

Lemma 6 (Posteriors at $a = 0$ histories). *Let λ be the congruent politician's probability of reaching $(1, 0)$.*

1. *If $\lambda > 0$, every non-congruent type choosing $a = 0$ uses $(1, 0)$, and $\mu_{10} > 0$.*
2. *If $\lambda = 0$, every on-path history with $a = 0$ has posterior zero. The distribution of shirkers across $(0, 0)$ and $(1, 0)$ is payoff irrelevant.*
3. *An off-path history with $a = 0$ is assigned posterior zero by D1.*

Proof. If $\lambda > 0$, the congruent type is present at $(1, 0)$. Therefore $\mu_{10} > 0$. Any on-path $(0, 0)$ can be generated only by non-congruent types by Lemma 4, so $\mu_{00} = 0$. Since consultation is costless and the rent is the same at the two status-quo histories, every shirker strictly prefers $(1, 0)$.

If $\lambda = 0$, neither history with $a = 0$ contains the congruent type, so every such history on the path of play has posterior zero. The two routes then give a shirker the same payoff. At an off-path history with $a = 0$, shirkers attain the smallest critical reelection probability because a sure relocation preserves their rent and their equilibrium reelection probability. The selected support contains only non-congruent types, and the posterior is therefore zero. $\square$

B.3 Posteriors and the cut-off

For $\lambda \in [0, \beta]$ and r , define

$$M_1(r; \lambda) \equiv \frac{\pi(1 - \lambda)}{\pi(1 - \lambda) + (1 - \pi)G(r)}, \quad (\text{B6})$$

$$M_0(r; \lambda) \equiv \frac{\pi\lambda}{\pi\lambda + (1 - \pi)[1 - G(r)]}. \quad (\text{B7})$$

At $\lambda = 0$, $M_1(r; 0) = \mu_1^B(r)$ and $M_0(r; 0) = 0$; at $\lambda = \beta$.

Lemma 7 (Bayes-consistent posteriors). *Suppose the non-congruent strategy has cutoff r and the congruent type reaches $(1, 0)$ with probability λ . Let σ_N be the probability that a*

non-congruent politician with rents lower than r consults. Then

$$\sigma_N = \frac{\sigma(1 - \beta\tau)}{1 - \lambda} \quad (\text{B8})$$

and consequently, each $a = 1$ on the path carries the common posterior $M_1(r; \lambda)$, while the history $(1, 0)$ carries $M_0(r; \lambda)$.

Proof. Suppose first that both $(0, 1)$ and $(1, 1)$ histories are on the path, so $0 < \sigma < 1$. The congruent type reaches $(0, 1)$ with probability $1 - \sigma$ and $(1, 1)$ with probability $\sigma[(1 - \beta) + \beta(1 - \tau)] = \sigma(1 - \beta\tau)$, and these sum to $1 - \lambda$. Applying Bayes' rule at each route separately,

$$\mu_{11} = \frac{\pi\sigma(1 - \beta\tau)}{\pi\sigma(1 - \beta\tau) + (1 - \pi)G(r)\sigma_N}, \quad \mu_{01} = \frac{\pi(1 - \sigma)}{\pi(1 - \sigma) + (1 - \pi)G(r)(1 - \sigma_N)}.$$

By Lemma 5, $\mu_{01} = \mu_{11}$, which holds if and only if the non-congruent-to-congruent ratio is the same at both routes:

$$\frac{1 - \sigma_N}{1 - \sigma} = \frac{\sigma_N}{\sigma(1 - \beta\tau)}.$$

This yields (B8). Substituting σ_N into μ_{01} ,

$$\mu_{01} = \frac{\pi(1 - \sigma)}{\pi(1 - \sigma) + (1 - \pi)G(r)\frac{1 - \sigma}{1 - \lambda}} = \frac{\pi(1 - \lambda)}{\pi(1 - \lambda) + (1 - \pi)G(r)} = M_1(r; \lambda),$$

after canceling $1 - \sigma > 0$ and multiplying through by $1 - \lambda$; the same substitution in μ_{11} cancels $\sigma(1 - \beta\tau) > 0$ and gives the identical expression.

If exactly one history with $a = 1$ is on the path of play, Lemma 5 places every mimicker there. When $\sigma = 1$ the congruent type reaches $(1, 1)$ with probability $1 - \beta\tau = 1 - \lambda$ and $\sigma_N = 1$, so Bayes' rule at $(1, 1)$ gives $M_1(r; \lambda)$. When $\sigma = 0$, the symmetric argument applies at $(0, 1)$, where the congruent type arrives with probability $1 = 1 - \lambda$.

For $a = 0$, Lemma 6 places every shirker at $(1, 0)$ when $\lambda > 0$, and the congruent type arrives there with probability λ , so Bayes' rule gives

$$\mu_{10} = \frac{\pi\lambda}{\pi\lambda + (1 - \pi)[1 - G(r)]} = M_0(r; \lambda),$$

while $\mu_{00} = 0$. □

Define the reputational return of mimicking $((0, 1)$ or $(1, 1))$ rather than shirking after consultation $(1, 0)$,

$$\Psi(r; \lambda) \equiv \nu[F(M_1(r; \lambda)) - F(M_0(r; \lambda))]. \quad (\text{B9})$$

Lemma 8 (The cutoff map). *For each $\lambda \in [0, \beta]$, non-congruent behavior is a cutoff rule and the equation $r = \Psi(r; \lambda)$ has a unique solution $\tilde{r}(\lambda) \in (\underline{r}, \bar{r})$. The map $\tilde{r}$ is continuous and strictly decreasing in λ , with*

$$\tilde{r}(0) = r^*, \quad \tilde{r}(\beta) \equiv \hat{r}, \quad \hat{r} < r^*. \quad (\text{B10})$$

Proof. The shirk-minus-mimic payoff $r + \nu F(M_0) - \nu F(M_1)$ is strictly increasing in r , so best responses are cutoffs. As r rises, $G(r)$ rises, M_1 strictly falls and M_0 strictly rises, so $\Psi(\cdot; \lambda)$ is strictly decreasing and $h(r; \lambda) \equiv r - \Psi(r; \lambda)$ is strictly increasing and continuous. At $r = \underline{r}$ we have $M_1 = 1$ and $M_0(\underline{r}; \lambda) = \frac{\pi\lambda}{\pi\lambda+1-\pi} < \pi$, so by Assumption 1, $\Psi(\underline{r}; \lambda) > \nu[F(1) - F(\pi)] > \underline{r}$ and $h(\underline{r}; \lambda) < 0$. At $r = \bar{r}$ we have $G = 1$: if $\lambda > 0$ then $M_0 = 1 \geq M_1$ and $\Psi \leq 0 < \bar{r}$; if $\lambda = 0$ then $\Psi(\bar{r}; 0) = \nu[F(\pi) - F(0)] < \bar{r}$ by Assumption 1. Either way $h(\bar{r}; \lambda) > 0$. The intermediate value theorem and strict monotonicity yield a unique interior solution $\tilde{r}(\lambda)$.

For an interior r , increasing λ strictly lowers M_1 and strictly raises M_0 . Hence $\Psi(r; \lambda)$ strictly decreases in λ . Continuity and strict monotonicity of h imply that its zero $\tilde{r}(\lambda)$ is continuous and strictly decreasing. At $\lambda = 0$, the fixed-point equation is the baseline cut-off; at $\lambda = \beta$, it is the full-compliance ($\sigma\tau = 1$) fixed-point equation. This establishes (B10). $\square$

B.4 Congruent incumbent's incentives

Lemma 9 (Compliance with a favorable recommendation). *In every D1-refined equilibrium, the congruent type chooses $a = 1$ after $m = 1$: That is, $\alpha_C(1) = 1$*

Proof. By Lemma 8 the cutoff is interior and $\nu[F(M_1) - F(M_0)] = \tilde{r}(\lambda) > 0$, so $a = 1$ strictly dominates the $(1, 0)$ on reputation. After $m = 1$, it also strictly dominates on policy by (B3) and the accompanying inequality. Hence $\alpha_C(1) = 1$, and with $\alpha_C(\emptyset) = 1$ the only congruent route to $(1, 0)$ is consultation followed by compliance to $m = 0$. $\square$

Lemma 10 (Incentive conditions). *Let $\lambda = \beta\sigma\tau$ and let the equilibrium cutoff be $\tilde{r}(\lambda)$. After $m = 0$, the congruent politician's payoff gain from following the recommendation rather than disregarding it is*

$$v - \tilde{r}(\lambda). \quad (\text{B11})$$

Before consultation, his expected payoff gain from consulting rather than choosing direct reform is

$$\beta\tau[v - \tilde{r}(\lambda)]. \quad (\text{B12})$$

Therefore:

1. if $0 < \tau < 1$, then $v = \tilde{r}(\lambda)$;

2. if $0 < \sigma < 1$ and $\tau > 0$, then $v = \tilde{r}(\lambda)$;
3. if $\tau = 0$, consultation does not change policy and σ is payoff irrelevant;
4. if $\sigma = \tau = 1$, compliance and consultation are optimal if and only if $v \geq \tilde{r}(\beta) = \hat{r}$.

Proof. By B2, compliance after $m = 0$ raises the expected policy payoff by v . It changes the public history from $(1, 1)$ to $(1, 0)$, and therefore lowers the reelection component by $\nu[F(M_1) - F(M_0)] = \tilde{r}(\lambda)$. This proves (B11).

Consultation differs from direct $a = 1$ only if the congruent incumbent complies with an unfavorable recommendation, an event of probability $\beta\tau$. The unconditional policy gain is $\tau(q - p) = \beta\tau v$, while the expected reputational loss is $\beta\tau\tilde{r}(\lambda)$. Their difference is (B12). The four implications are the usual indifference and weak-optimality conditions at the two information sets. $\square$

Lemma 11 (Off-path beliefs). *The off-path beliefs needed for equilibrium verification are as follows.*

1. At every off-path history with $a = 0$, the posterior is zero.
2. In a non-compliance equilibrium, an off-path history with $a = 1$ is a D1 tie between the congruent type and the equilibrium mimickers and is assigned $\mu_1^B(r^*)$.
3. In a partial-compliance equilibrium with $\sigma = 1$, direct $a = 1$ is a D1 tie between the congruent type and the mimickers and is assigned $M_1(\tilde{r}(\lambda); \lambda)$.
4. In the full-compliance equilibrium, direct $a = 1$ is assigned posterior zero when $v > \hat{r}$ and the tie-preserving posterior $M_1(\hat{r}; \beta)$ when $v = \hat{r}$.

Proof. (1) is Lemma 6. (2) An off-path history with $a = 1$ delivers the same policy payoff as the on-path one with $a = 1$ for the congruent type and for every mimicker, so all tie at the critical value $F(\mu_1^B(r^*))$, while a shirker has $\frac{r}{\nu} + F(0) > F(\mu_1^B(r^*))$; the tie-preserving belief is $\mu_1^B(r^*)$. (3) The condition $v = \tilde{r}(\lambda)$ makes the congruent equilibrium payoff $p + \nu F(M_1)$, so his critical value at direct $a = 1$ is $F(M_1)$, exactly that of each mimicker, while shirkers are strictly higher. (4) In full compliance $U_C = q + \nu[(1 - \beta)F(M_1) + \beta F(M_0)]$ with $M_1 = M_1(\hat{r}; \beta)$ and $M_0 = M_0(\hat{r}; \beta)$; using $\hat{r} = \nu[F(M_1) - F(M_0)]$,

$$x_C(0, 1) - x_M(0, 1) = \frac{q - p - \beta\hat{r}}{\nu} = \frac{\beta(v - \hat{r})}{\nu},$$

so only mimickers minimize when $v > \hat{r}$, and the congruent type joins them at $v = \hat{r}$. $\square$

C Proof of Proposition 2

Proof of Proposition 2. By Lemmas 4 and 9 the congruent type chooses $a = 1$ whenever he does not consult and after $m = 1$, so his only route to $a = 0$ is $(1, 0)$, reached with probability $\lambda = \beta\sigma\tau$. Lemmas 5–6 locate non-congruent mimickers at the $a = 1$ histories alongside the congruent type and non-congruent shirkers at $(1, 0)$ when $\lambda > 0$, and Lemma 8 delivers the unique interior cutoff $\tilde{r}(\lambda)$, continuous and strictly decreasing with $\tilde{r}(0) = r^*$. Every D1-refined equilibrium is therefore indexed by λ , and there are two cases.

Type I ($\lambda = 0$, equivalently $\sigma\tau = 0$). The cutoff is $\tilde{r}(0) = r^*$, every on-path history with $a = 1$ has posterior $\mu_1^B(r^*)$, and every history with $a = 0$ has posterior zero: on the path by Lemma 6 and off the path by Lemma 11(1)–(2). The congruent type chooses $a = 1$ with probability one whether or not he consults: under avoidance $\sigma = 0$; under pandering $\tau = 0$ and consultation is reputationally inert, so σ is indeterminate. A non-congruent type with $r < r^*$ mimics, and one with $r \geq r^*$ chooses $a = 0$, and by Lemma 6 he is indifferent between doing so directly and after consulting.

Type II ($\lambda > 0$, equivalently $\sigma, \tau > 0$). The cutoff is $\tilde{r}(\lambda) < r^*$ by strict monotonicity. Every on-path history with $a = 1$ has posterior $M_1(\tilde{r}(\lambda); \lambda)$, while $\mu_{10} = M_0(\tilde{r}(\lambda); \lambda) > 0$ and $\mu_{00} = 0$, with the $\sigma = 1$ tie handled by Lemma 11(3)–(4). The congruent type consults with probability $\sigma > 0$, chooses $a = 1$ when he does not consult and after $m = 1$, and obeys $m = 0$ with probability $\tau > 0$. By Lemma 6, every non-congruent type above the cutoff consults and then chooses $a = 0$ with certainty, since $(1, 0)$ strictly dominates $(0, 0)$ on reputation at equal rent.

It remains to verify that no omitted behavioral deviation is profitable. For non-congruent types, Lemma 3 reduces signal-contingent deviations to convex combinations governed by the cutoff, and relocation across histories cannot improve the posterior by Lemmas 5, 6 and 11. For the congruent type, direct $a = 1$ and compliance to $m = 1$ are Lemmas 4 and 9, and the remaining consultation and compliance following an unfavorable recommendation's decisions are characterized by Lemma 10. The two classes therefore exhaust the D1-refined equilibrium set. $\square$

C.1 Proof of Corollary 1

Proof. If $\kappa = 0$ then $\lambda = 0$ and $\tilde{r}(0) = r^*$, and every history with $a = 0$ has posterior zero by Lemma 6, so both margins reproduce Proposition 1.

If $\kappa > 0$ then $\lambda = \beta\kappa > 0$, and Lemma 8 gives $\tilde{r}(\lambda) < r^*$ with $\tilde{r}$ strictly decreasing, so the displaced mass $G(r^*) - G(\tilde{r}(\lambda))$ is positive and strictly increasing in λ , hence in κ . By

Lemma 7,

$$\mu_{10} = M_0(\tilde{r}(\lambda); \lambda) = \frac{\pi\lambda}{\pi\lambda + (1 - \pi)[1 - G(\tilde{r}(\lambda))]} > 0,$$

so $F(\mu_{10}) > F(0)$. Raising λ increases the numerator directly and lowers $G(\tilde{r}(\lambda))$, which increases $1 - G(\tilde{r}(\lambda))$; since the numerator rises proportionally faster than the denominator, M_0 is strictly increasing in λ , and so is $F(\mu_{10})$. $\square$

D Proof of Proposition 3

Proof of Proposition 3. Recall $v = (q - p)/\beta$ from (B4) and $\hat{r} = \tilde{r}(\beta)$ from (B10).

(a) *Full compliance.* Set $\sigma = \tau = 1$, so $\kappa = 1$ and $\lambda = \beta$. By Lemma 10(4), consulting and obeying $m = 0$ are optimal for the congruent type if and only if $v \geq \hat{r}$. Under Lemma 11(4) the deviation to direct $a = 1$ is unprofitable: when $v > \hat{r}$ it draws posterior zero, and when $v = \hat{r}$ it draws $M_1(\hat{r}; \beta)$ and leaves the congruent type exactly indifferent. Non-congruent deviations are ruled out by the cutoff and Lemma 3. Hence, a full-compliance equilibrium exists if and only if $v \geq \hat{r}$, and the inequality is strict except at the boundary $v = \hat{r}$, where the profile also arises as the $\kappa \rightarrow 1$ limit of part (b).

(b) *Partial compliance.* Suppose $0 < \kappa < 1$, so $0 < \lambda < \beta$ and at least one of σ, τ is interior. By Lemma 10(1)–(2), either interim indifference after $m = 0$ or indifference at the consultation stage requires $v = \tilde{r}(\lambda) = \tilde{r}(\beta\kappa)$. Conversely, if $v = \tilde{r}(\beta\kappa)$ then both gains in Lemma 10 vanish, so every pair (σ, τ) with $\sigma\tau = \kappa$ is sequentially rational; beliefs are given by Lemma 7 and, at $\sigma = 1$, by Lemma 11(3). Since $\tilde{r}$ is continuous and strictly decreasing from r^* to $\hat{r}$, such a κ exists and is unique precisely when $\hat{r} < v < r^*$, in which case $\kappa^* = \tilde{r}^{-1}(v)/\beta$. $\square$

Remark 1 (The equilibrium correspondence). Combining Proposition 3 with the Type-I conditions of Proposition 2, a non-compliance equilibrium exists if and only if $v \leq r^*$: at $\lambda = 0$ the cost of compliance is $\tilde{r}(0) = r^*$, deviating to consult-and-obey gains $\beta(v - r^*)$ and pandering gains $v - r^*$, while the tie-preserving beliefs of Lemma 11(2) make switches between direct $a = 1$ and inert consultation neutral. The full set of effective compliance rates is therefore

region	compliance rates κ
$v < \hat{r}$	0
$v = \hat{r}$	0, 1
$\hat{r} < v < r^*$	0, κ^* , 1
$v = r^*$	0, 1
$v > r^*$	1.

When $v > r^*$ full compliance is the unique D1-refined equilibrium. When $\hat{r} < v < r^*$ the interior rate κ^* is unique but its decomposition is not: the implementation set is $\{(\sigma, \tau) : \sigma \in [\kappa^*, 1], \tau = \kappa^*/\sigma\}$, and all such pairs are outcome-equivalent, since the posteriors, the cutoff, and the congruent type's expected accuracy depend on (σ, τ) only through $\lambda = \beta\sigma\tau$.

E Proof of Proposition 4

Lemma 12 (Policy accuracy by type). *With effective compliance κ and cutoff $\tilde{r} = \tilde{r}(\beta\kappa)$, the congruent type's expected accuracy is $p + \kappa(q - p)$; a mimicker is correct with probability p and a shirker with probability $1 - p$.*

Proof. Absent compliance the congruent type chooses $a = 1$ and is correct with probability p . An unfavorable recommendation is obtained and obeyed with probability $\beta\kappa$, and each such event raises his accuracy by $v = (q - p)/\beta$ from (B4), for a total gain of $\beta\kappa v = \kappa(q - p)$. Non-congruent policies are state independent: mimickers choose $a = 1$ and match with probability p , shirkers choose $a = 0$ and match with probability $1 - p$. $\square$

Proof of Proposition 4. By Lemma 12, welfare in an equilibrium with compliance κ and cutoff $\tilde{r} = \tilde{r}(\beta\kappa)$ is

$$W(\kappa) = \pi[p + \kappa(q - p)] + (1 - \pi)[G(\tilde{r})p + (1 - G(\tilde{r}))(1 - p)],$$

and, setting $\kappa = 0$ and $\tilde{r} = r^*$, benchmark welfare is

$$W_B = \pi p + (1 - \pi)[G(r^*)p + (1 - G(r^*))(1 - p)].$$

Subtracting and collecting the non-congruent terms, each type between $\tilde{r}$ and r^* switches from accuracy p to accuracy $1 - p$, a loss of $2p - 1$ per unit of mass, so

$$W(\kappa) - W_B = \underbrace{\pi\kappa(q - p)}_{\text{expertise gain}} - \underbrace{(1 - \pi)[G(r^*) - G(\tilde{r}(\beta\kappa))]}_{\text{discipline loss}}(2p - 1),$$

which is equation (9) of the main text. Advice raises welfare if and only if the first term weakly exceeds the second; at full compliance this is $\pi(q - p) \geq (1 - \pi)[G(r^*) - G(\hat{r})](2p - 1)$. $\square$

F Comparative statics

The claims in Sections 4.3 and 5 of the main text compare the benchmark with the full-compliance equilibrium, restricting attention to parameter changes that preserve existence. Write $D^B(r, \pi) = F(\mu_1^B(r)) - F(0)$ and $D^A(r, \pi, q) = F(M_1(r; \beta)) - F(M_0(r; \beta))$, so that $r^* = \nu D^B(r^*, \pi)$ and $\hat{r} = \nu D^A(\hat{r}, \pi, q)$.

Lemma 13 (Implicit-function denominators). *With $\Gamma_B = 1 - \nu D_r^B(r^*, \pi)$ and $\Gamma_A = 1 - \nu D_r^A(\hat{r}, \pi, q)$, one has $D_r^B < 0$, $D_r^A < 0$, and hence $\Gamma_B, \Gamma_A > 0$.*

Proof. μ_1^B is strictly decreasing in r , so $D_r^B < 0$. In $M_1(\cdot; \beta)$ and $M_0(\cdot; \beta)$ the former is strictly decreasing and the latter strictly increasing in r , and $F' > 0$, so $D_r^A < 0$. $\square$

Proposition 5 (Advisor accuracy). *$\partial r^*/\partial q = 0$ and $\partial \hat{r}/\partial q > 0$. Consequently, the discipline loss $G(r^*) - G(\hat{r})$ is strictly decreasing in q and the welfare gap $W(1) - W_B$ is strictly increasing in q .*

Proof. The benchmark cutoff map does not involve q , so $\partial r^*/\partial q = 0$. Since $\partial \beta/\partial q = 1 - 2p < 0$, differentiating (B6) and (B7) at $\lambda = \beta$ gives

$$\frac{\partial M_1}{\partial q} = (2p-1) \frac{\pi(1-\pi)G(r)}{[\pi(1-\beta) + (1-\pi)G(r)]^2} > 0, \quad \frac{\partial M_0}{\partial q} = -(2p-1) \frac{\pi(1-\pi)[1-G(r)]}{[\pi\beta + (1-\pi)(1-G(r))]^2} < 0,$$

so $D_q^A > 0$ and $\partial \hat{r}/\partial q = \nu D_q^A/\Gamma_A > 0$ by Lemma 13. Hence $G(r^*) - G(\hat{r})$ falls in q , and by Proposition 4

$$\frac{\partial [W(1) - W_B]}{\partial q} = \pi + (1-\pi)(2p-1)g(\hat{r}) \frac{\partial \hat{r}}{\partial q} > 0. \quad \square$$

Proposition 6 (Value of office). *$\partial r^*/\partial \nu > 0$ and $\partial \hat{r}/\partial \nu > 0$; the signs of $\partial [G(r^*) - G(\hat{r})]/\partial \nu$ and of the corresponding welfare derivative are ambiguous.*

Proof. Implicit differentiation gives $\partial r^*/\partial \nu = r^*/(\nu \Gamma_B) > 0$ and $\partial \hat{r}/\partial \nu = \hat{r}/(\nu \Gamma_A) > 0$. But

$$\frac{\partial [G(r^*) - G(\hat{r})]}{\partial \nu} = g(r^*) \frac{\partial r^*}{\partial \nu} - g(\hat{r}) \frac{\partial \hat{r}}{\partial \nu}$$

compares two positive terms and with an ambiguous sign, so is $\partial [W(1) - W_B]/\partial \nu$. $\square$

Proposition 7 (Prior congruence). *$\partial r^*/\partial \pi > 0$; the sign of $\partial \hat{r}/\partial \pi$ is ambiguous, and so is that of the corresponding welfare derivative.*

Proof. Since $\partial\mu_1^B/\partial\pi = G(r)/[\pi + (1 - \pi)G(r)]^2 > 0$, one has $D_\pi^B > 0$ and $\partial r^*/\partial\pi = \nu D_\pi^B/\Gamma_B > 0$. At $\lambda = \beta$ both $\partial M_1/\partial\pi > 0$ and $\partial M_0/\partial\pi > 0$, so

$$\frac{\partial \hat{r}}{\partial \pi} = \frac{\nu}{\Gamma_A} \left[F'(M_1) \frac{(1 - \beta)G(\hat{r})}{[\pi(1 - \beta) + (1 - \pi)G(\hat{r})]^2} - F'(M_0) \frac{\beta[1 - G(\hat{r})]}{[\pi\beta + (1 - \pi)(1 - G(\hat{r}))]^2} \right]$$

is a difference of positive terms with ambiguous sign. Consequently,

$$\frac{\partial[W(1) - W_B]}{\partial\pi} = (q - p) + [G(r^*) - G(\hat{r})](2p - 1) - (1 - \pi)(2p - 1) \left[g(r^*) \frac{\partial r^*}{\partial\pi} - g(\hat{r}) \frac{\partial \hat{r}}{\partial\pi} \right]$$

is ambiguous. □