

Governing Through Division*

Giovanna M. Invernizzi[†] Federico Trombetta[‡]

May 27, 2025

Abstract

We study how party factionalization affects policy-making. We develop an accountability model where the incumbent party leader can implement reforms but faces internal factions that can reduce reform effectiveness through dissent. The leader is aligned with voter preferences, while factions have competing policy goals but share the leader's re-election objectives. Voters are uncertain about whether a reform is needed and how divided the party is, and value strong leaders who can effectively implement reforms. In equilibrium, dissent creates the informational environment that facilitates over-reform: precisely because they can overcome factional resistance, strong leaders reform even when maintaining the status quo would be optimal. In turn, strong factionalization can improve policymaking by removing the temptation to over-reform: weak leaders behave well not despite their weakness, but because of it. The welfare effects are non-monotonic: intermediate factionalization optimally balances the trade-off between reform effectiveness and equilibrium incentives to choose the correct policy. We also show that active dissent can yield higher voter welfare than strategic silence by factions.

*We thank Tim Besley, Massimo Bordignon, Peter Buisseret, Salvatore Di Salvio, Livio Di Lonardo, Gabriele Gratton, Helios Herrera, Niall Hughes, Barton Lee, Massimo Morelli, Salvatore Nunnari, Carlo Prato, Federico Vaccari, and Richard van Weelden for their insightful suggestions. We thank participants at King's College London, Bocconi PE Breakfast, UniCatt Pol Econ WIP seminar, 3rd IE-Rochester Political Economy workshop for their comments.

[†]Department of Social and Political Sciences, Bocconi University. Email: giovanna.invernizzi@unibocconi.it.

[‡]Department of International Economics, Institutions and Development, Faculty of Political and Social Sciences, Università Cattolica del Sacro Cuore, Largo Gemelli 1, 20123 Milano (Italy).

1. Introduction

Political parties are complex organizations characterized by internal divisions and competing interests. These divisions often manifest through factions—groups within the party that maintain distinct ideological positions, identities, and policy goals. The role of factions in democratic governance has been debated since the inception of modern democracy, with the Founding Fathers expressing concerns about their potential to undermine effective governance ([Madison, 1787](#)).

This historical skepticism has largely persisted in contemporary scholarship, where factions are predominantly viewed as impediments to effective policy-making ([Rohde, 1991](#); [Tsebelis, 2002](#)). This perspective stems from factions’ ability to obstruct legislative bargaining ([Wawro and Schickler, 2013](#)), their tendency to increase coordination costs in legislative decision-making ([Cox and McCubbins, 2007](#)), and their diversion of resources through internal rent-seeking competition ([Persico et al., 2011](#)), thereby introducing frictions into the policy-making process. In turn, these frictions are taken into account by voters, who prefer cohesive parties over divided ones, often penalizing those perceived as internally divided ([Greene and Haber, 2015](#); [Lehrer et al., 2024](#)).

Yet these same frictions that appear detrimental may offer unexpected benefits for democratic accountability. This paper challenges the conventional wisdom that party factionalization necessarily reduces voter welfare. While we begin from the widely-accepted premise that factions hamper policy implementation, we demonstrate that these impediments can serve a beneficial role by mitigating moral hazard problems between governing parties and voters. Our key insight is that the strategic interaction between party leaders and their internal factions creates a disciplining effect on policy-making. Strong leaders — those facing weak factions that cannot effectively oppose them — may implement unnecessary reforms to signal their control over the party. They do so even when aligned with voters’ preferences, and voters reward them electorally despite recognizing these reforms as sub-optimal. Conversely, leaders facing factions capable of meaningful dissent are not able to over-reform, hence do not choose unnecessary reforms. These opposing forces create conditions where voter welfare can increase with factionalization.

The tendency of party leaders to implement ambitious and potentially unnecessary reforms to demonstrate their strength is evident in several historical cases. A striking example is President Trump’s 2018 tariff policies on steel and aluminum imports. Despite facing significant opposition from Republican lawmakers and business interests within his own party, Trump was able to impose 25% tariffs on steel and 10% tariffs on aluminum imports, overriding internal dissent and

starting a trade war with China. Congressional Republicans expressed their opposition through various means—from Senator Jeff Flake’s proposal to introduce legislation nullifying the tariffs to a letter signed by over 100 Republican lawmakers imploring the president to reconsider (CBS News, 2018).

Interestingly, Autor et al. (2024) find that, while these tariffs had relatively small employment effects in protected sectors and negative impacts in sectors hit by counter-tariffs, they generated substantial political benefits for Trump and the Republican party. Their research reveals that “local exposure to the trade war appears to have benefited the Republican party even in regions where the combined effect of tariffs and subsidies predicts no employment gain or even a modest employment loss” (p. 25). This political benefit manifested in increased support for tariffs, stronger identification with the Republican party, and higher Republican vote share in regions with many newly protected industries.

While Autor et al. (2024) suggest explanations based on voter misinformation or appreciation of Trump’s intentions, our model offers a fully rational alternative: voters responded positively to the demonstration of leadership strength through the successful implementation of controversial reforms. In our framework, this occurs in equilibrium precisely because strong leaders can signal their capacity to overcome internal opposition by implementing reforms — even suboptimal ones — that weaker leaders cannot execute. Importantly, our explanation requires neither voter inattention nor expressive voting motivations: it demonstrates how electoral benefits can emerge from policy choices that voters rationally recognize as inefficient. This aligns with the observed political success of Trump’s tariff policies, which served as a credible signal of his leadership effectiveness despite their questionable economic merits.

To formalize this dynamic, we develop an accountability model where the leader of an incumbent party must decide whether to implement reforms while contending with internal factions that can potentially obstruct their implementation. The leader observes the state of the world and knows whether a reform is truly needed. Voters, as well as the party leader, suffer a policy cost when the policy outcome does not match the state of the world. In other words, voters want a reform to be implemented only when needed, and the party leader is perfectly aligned with voters’ preferences. However, factions hold different policy preferences: while they agree with the leader when no reform is needed, they prefer a different direction when reforms are necessary. When factions choose to dissent, they can obstruct implementation, tilting the policy outcome away from the leader’s proposal. The effectiveness of this obstruction depends on the faction’s

strength relative to the leader (a measure of their relative bargaining power). Empirically, such factional opposition manifests through various obstructionist tactics—from procedural delays in legislative chambers to bureaucratic resistance during policy implementation—all of which can effectively sabotage the leader’s reform agenda.

Voters face uncertainty along two dimensions: they do not observe whether reforms are truly needed, and they are uncertain about the degree of party factionalization. The latter can also be interpreted as uncertainty over the leader’s strength, which we model as the leader’s ability to overcome party internal opposition. Hence, the more factionalized a party, the weaker its leader. This uncertainty reflects the observation that citizens possess incomplete knowledge of party internal divisions (Greene and Haber, 2015) and the fact that party leaders strategically manage public perceptions through unity displays, further obscuring the true extent of internal opposition. We assume that in the absence of reform voters do not (directly) learn the incumbent’s strength. In contrast, successfully implementing a reform can reveal how strong the party leader is (and how factionalized the party is).

Our analysis reveals several results. First, we ask when factions dissent in equilibrium. When party leaders can retaliate against dissenting factions by reducing their resource allocation, factions face a trade-off between immediate policy influence and future rents. We show that factions dissent in equilibrium when their share of party resources is sufficiently insulated from leadership discretion—a feature more common in institutionalized parties with transparent allocation mechanisms. This generates an empirical prediction: we should observe more active factional dissent in parties where internal resource distribution follows objective rules rather than leadership discretion.

Second, we show that strong party leaders—those whose factions cannot effectively obstruct—may implement unnecessary reforms to demonstrate their control. This “Over-Reform” equilibrium emerges because voters, who value effective governance, interpret successful reform implementation as a signal of party cohesion. Consequently, even when maintaining the status quo would be optimal, strong leaders may initiate reforms to distinguish themselves from weaker leaders who face substantial internal opposition. This creates a perverse incentive where the most capable leaders engage in inefficient policy-making precisely because of their ability to overcome internal opposition. This equilibrium arises even when voters can perfectly observe the consequences of a given policy choice.

The welfare implications of this mechanism challenge the conventional view that having a strong leadership unambiguously benefits voters. We demonstrate that voter welfare can be non-monotonic (reverse U-shaped) in the expected level of party factionalization when uncertainty exists about faction strength. While having a strong leadership (or weak factions) in the incumbent party enhances implementation of necessary reforms, it also increases the likelihood of unnecessary reforms undertaken for signalling purposes. Since welfare gains from better policy implementation diminish as the probability of electing strong leaders increases, welfare reaches its maximum at intermediate levels of expected leadership strength. At this optimal point, the probability of factions strong enough to prevent excessive reforms balances with the likelihood of effective necessary policy changes. This finding suggests that strong party factions, rather than being purely detrimental, can serve as beneficial constraints on leaders' tendency to over-reform—effectively functioning as commitment devices that make political restraint credible.

Interestingly, we find that voters may sometimes prefer active factions despite their misalignment with voter preferences. When factions strategically choose not to dissent—typically when dissent would jeopardize their future resource allocation within the party—they enable both necessary and unnecessary reforms to pass unobstructed. Under certain conditions, specifically when strong leaders are common and reforms are rarely needed, active factional dissent that constrains all reforms can yield higher voter welfare than passive factions that allow over-reform. This counterintuitive result emerges because the constraint on unnecessary reforms, and the additional information brought to the voter by the policy outcome under factional dissent, outweighs the cost of obstructing beneficial ones.

Besides the over-reform equilibrium described above, there also exists an equilibrium where strong and weak leaders pool on always implementing the correct reform, which is standard in this class of models. We refer to this as the “Full Discipline equilibrium.” While intuitively one might expect the Full Discipline equilibrium to dominate the Over-Reform one in terms of voter welfare, we demonstrate that the latter can actually yield higher voter welfare under specific conditions. This occurs because the informational benefits of separating incumbent types can outweigh the policy costs of occasional unnecessary reforms, particularly when the probability of a strong incumbent is sufficiently low and the effectiveness gap between strong and weak incumbents is substantial.

Our findings contribute to ongoing debates about the role of party unity in democratic governance by highlighting a novel mechanism through which internal divisions might unexpectedly

benefit voters. While most of the existing literature has emphasized the negative effects of factionalization on policy implementation, our model reveals its potential disciplining effect on strategic policy choices. This perspective offers new insights for interpreting empirical patterns of reform activity across political systems with varying degrees of party cohesion. Moreover, our theoretical framework provides a foundation for future empirical investigations into the relationship between party structure, reform initiatives, and electoral accountability—particularly in contexts where signalling considerations might drive policy decisions.

The remainder of the paper proceeds as follows. Section 2 discusses our contribution to the existing literature. Section 3 introduces our model, while Section 4 present the central results. Section 5 concludes.

2. Related Literature

Our paper begins with the premise that parties are not monolithic actors but rather consist of different factions with competing interests and goals. While political scientists have long acknowledged the existence of factions, recent literature has increasingly recognized their critical role in shaping various political outcomes. These include party nomination processes (Caillaud and Tirole, 2002; Crutzen et al., 2010), intra-party power sharing (Persico et al., 2011; Invernizzi and Prato, 2024), internal party conflict (Invernizzi, 2023; Izzo, 2024), and legislative decision-making (Cox and McCubbins, 2007; Wawro and Schickler, 2013). Within this literature, factions are predominantly viewed as impediments to effective governance—creating coordination problems, increasing transaction costs, and introducing friction into policy implementation.

Our contribution challenges this conventional wisdom by demonstrating that factional opposition can sometimes benefit voters through a novel mechanism: by constraining party leaders from implementing unnecessary reforms for signalling purposes. Dewan and Squintani (2016) provide another exception to the predominantly negative view of factions, by showing how factions can improve welfare by enhancing internal party deliberations. Rather than information aggregation, we show that factions serve as a beneficial constraint on strategic policy choices, especially when electoral incentives might otherwise lead to inefficient policy-making.

To isolate the effects of factional constraints on policy-making, we deliberately model parties as fixed entities with established internal divisions. In reality, party composition is more fluid, with factions sometimes choosing to splinter and form new political entities (Invernizzi and Izzo, 2025). Such exit dynamics can be conceptualized as occurring at an earlier stage than our

model—we focus instead on the subsequent phase where factions have already committed to remain within the party but express dissent internally. This framing allows us to address how factional dissent shapes the strategic policy choices of party leadership. Close to our work is [Izzo \(2024\)](#), who examines how factions strategically employ dissent as a form of position-taking aimed at obtaining ideological benefits, in a model of policy experimentation. We instead analyze how dissent affects leadership behavior through its effects on policy implementation. This distinction generates a novel trade-off: factions can improve welfare by preventing excessive reforms driven by signalling motives (rather than by affecting voters’ learning about their own preferences), while negatively affecting policy effectiveness in general.

Political agency models provide the natural analytical framework to study how factions affect policy-making. This literature typically addresses two types of uncertainty. First, models examining uncertainty over politicians’ preferences or bias ([Acemoglu et al., 2013](#); [Kartik and Van Weelden, 2019](#); [Merzoni and Trombetta, 2022](#); [Lodato et al., 2024](#)) show how ‘good’ politicians may take inefficient actions to signal they are not biased. Close to us, [Lodato et al. \(2024\)](#) use a multi-level model of accountability, with (possibly biased) politicians and bureaucrats. Our mechanism differs fundamentally as inefficient behavior emerges without any underlying bias—it is driven purely by electoral incentives and asymmetric information about the relative strength of party leadership and factions. Second, research on uncertainty over politician competence ([Canes-Wrone et al., 2001](#); [Ashworth and Shotts, 2010](#); [Fox and Stephenson, 2011](#)) typically finds that high-ability politicians make better decisions while low-ability ones make mistakes or engage in inefficient behavior. In contrast, we show that the most capable politicians (those leading united parties) engage in inefficient behavior precisely because of their ability to overcome internal opposition. This connects our work to the literature on anti-herding ([Levy, 2004](#)). Our model also relates to the literature on accountability and “veto players” of different forms ([Buisseret, 2016](#); [Fox and Stephenson, 2011](#); [Fox and Van Weelden, 2010](#); [Caillaud and Tirole, 1999](#); [Hirsch and Kastellec, 2022](#)). Also in our case, the actions of the second player affect the informativeness of the implemented policy, but it does so asymmetrically across policies and types. Moreover, we study the role of a particular type of “additional player,” sharing the same objective as the party leadership but with different ideological positions.

Beyond political agency, our work relates to broader literature on institutional constraints and reform production. Various studies show how constraints—whether in voter information processing ([Prat, 2005](#); [Fox and Van Weelden, 2012](#); [Ashworth and Bueno De Mesquita, 2014](#)),

organizational communication (Che et al., 2013), or delegation (Dessein, 2002)—can restrict opportunistic behavior and improve outcomes. Our contribution shows that such constraints can be beneficial even without politicians’ bias or incompetence. Additionally, we connect to research on the over-production of reforms, where electoral incentives lead to excessive policy-making under rational inattention (Prato and Wolton, 2018), opponent campaigns (Dewan and Hortala-Vallve, 2019), or bureaucratic influence (Gratton et al., 2021). Our over-reform equilibrium is also similar to Shaver (2025), where the politician’s desire to signal his quality motivates inefficient policy choices. Our model highlights how the interaction between internal party politics and electoral incentives creates similar inefficiencies (and novel welfare implications) through an unexplored mechanism.

Our paper contributes to the policy experimentation literature by examining how asymmetric informativeness of different policies affects political behavior. Unlike studies focused on the policy environment (Blumenthal, 2024), we examine information about the politician’s type. While earlier models show how reputational concerns can lead politicians to choose more informative policies—either to signal competence (Majumdar and Mukand, 2004; Fu and Li, 2014) or learn about voter preferences (Bils and Izzo, 2025)—our mechanism operates differently in two key respects. First, in our equilibrium, it is specifically the high-ability politicians who inefficiently implement reforms to separate themselves from low-ability ones. Second, we model how internal party factions endogenously constrain signaling behavior.¹

Our result on inefficient effort allocation relates to multi-tasking models (Daley and Snowberg, 2011; Buisseret and Prato, 2016). Daley and Snowberg (2011) find that high types allocate effort to informative but wasteful tasks like fundraising. While our focus is on factions rather than campaign finance, we similarly find a separating equilibrium with high types choosing suboptimal behavior. However, in our case, “low” types may actually be more beneficial from the voter’s perspective, yielding important welfare implications.

¹Our framework incorporates asymmetric informativeness of different policies—implementing a reform reveals more about party strength than maintaining the status quo. Differently from Daley and Snowberg (2011) and Blumenthal (2024), whose focus is on the policy environment, we focus on information about the politician’s type. In this respect, we are close to Fu and Li (2014) in having an equilibrium where high types choose suboptimal policies. However, differently from their model (where policies reveal differential information conditional on observing also the outcome), in our equilibrium low types behave optimally. Moreover, they do not focus on (endogenous) factions, but rather on the welfare effect of different institutional environments.

3. The Model

Consider a two-period game with the following players: an incumbent party leader (henceforth incumbent, *he*), a faction, a representative voter (*she*), and a non-strategic challenger. In each period, there is a binary state of the world ($\omega_t \in \{0, 1\}$): when $\omega_t = 1$, implementing a reform is optimal for the voter, while $\omega_t = 0$ indicates that maintaining the status quo is optimal. The probability that reforming is optimal is $\Pr(\omega_t = 1) = \pi \in (0, 1)$.

In each period t , the incumbent chooses whether to implement a reform ($x_t = 1$) or not ($x_t = 0$). The faction can obstruct policymaking by dissenting: formally, the faction chooses $d_t \in \{0, 1\}$, where $d_t = 0$ corresponds to no dissent and $d_t = 1$ to dissent.²

We assume that the implemented policy is a function of the policy decision by the incumbent and the internal opposition that he faces within the party. Formally, the implemented policy in period t (defined as $\tilde{x}_t$) is

$$\tilde{x}_t = \begin{cases} x_t & \text{if } d_t = 0 \\ \phi_I x_t & \text{if } d_t = 1, \end{cases}$$

where $\phi_I \in \{\phi_L, \phi_H\}$ measures I 's relative strength vis-à-vis the faction, or how “flexible” the incumbent is to implement his chosen policy ($0 < \phi_L < \phi_H < 1$). In other words, we are assuming that the strong faction's dissent is more effective than the weak faction's. This is also a parsimonious way to capture the relative bargaining power between leadership and faction.

The incumbent and the faction have complete information about all relevant parameters. The voter, however, has uncertainty about the incumbent leader's strength: she believes that $\phi_I = \phi_H$ with probability $\gamma \in (0, 1)$, and $\phi_I = \phi_L$ otherwise. For simplicity, we assume that the challenger's strength parameter ϕ_C is distributed according to the same distribution as the incumbent's, and relax this assumption in Appendix C. While the voter does not know the state of the world, she observes the implemented policy $\tilde{x}_t$.³ Note that $\tilde{x}_1 = \phi_I$ perfectly reveals the incumbent's type.⁴

²Appendix D shows that results are robust to a continuous choice, where the faction decides “how much” to dissent.

³Our results do not rely on this assumption: the voter could also observe the state of the world and the choice of the incumbent. We further discuss this point below.

⁴We relax this assumption in Appendix B where we consider probabilistic revelation.

The voter simply wants a reform choice that matches the state of the world, and the incumbent is aligned with the voter's preferences. Because factions induce frictions in the policy-making process (and the incumbent is aligned), the voter prefers a strong incumbent to a weak one (when a reform is needed), because the former is more effective in implementing reforms. Formally, the voter's per-period payoff is:

$$u_t^v = -f(\tilde{x}_t, \omega_t),$$

where $f(\tilde{x}_t, \omega_t)$ is a single-peaked continuous function such that $f(\omega_t, \omega_t) = 0$, $f(\tilde{x}_t, \omega_t) > 0$ for all $\tilde{x}_t \neq \omega_t$, and $f(\tilde{x}_t, \omega_t)$ is strictly increasing in $|\tilde{x}_t - \omega_t|$ for any ω_t and $\tilde{x}_t$. One example is $f(\tilde{x}_t, \omega_t) = (\omega_t - \tilde{x}_t)^2$.

The payoff of the incumbent leadership when in office is:

$$v_t^I = (1 - \beta_t)R - f(\tilde{x}_t, \omega_t),$$

where R represents the rents from holding office and $1 - \beta_t$ is the leader's share of it. When out of office, the politician's payoff is normalized to zero.

Factions constitute an integral part of the party—sharing the party's electoral goals—yet maintain distinct interests from the leadership. This distinction manifests in two dimensions: resource allocation and ideological preferences. Formally, the faction obtains a lower share of office rents, which we denote by $\beta_t \in (0, 1/2)$. In terms of policy preferences, while the incumbent leadership always wants to match the state of the world, the faction only wants to match when $\omega = 0$. Instead, when $\omega = 1$, the policy preferred by the faction is $\hat{x}^F = -1$.⁵ That is, while at the status quo leader and faction agree that inaction is optimal, when there is a need for reform they disagree on what the optimal reform is. We can then write the payoff of the faction at time t as:

$$v_t^F = \beta_t R - g(\tilde{x}_t, \omega_t),$$

where $g(\tilde{x}_t, \omega_t)$ is a continuous single-peaked function such that $g(\omega_t, -\omega_t) = 0$, $g(\tilde{x}_t, \omega_t) > 0$ for all $\tilde{x}_t \neq \omega_t$, and $g(\tilde{x}_t, \omega_t)$ is strictly increasing in $|\tilde{x}_t + \omega_t|$ for any ω_t and $\tilde{x}_t$. One example is $g(\tilde{x}_t, \omega_t) = (\omega_t + \tilde{x}_t)^2$.

⁵Our results do not rely on this assumption. We discuss this point more in depth below.

We assume that $\beta_1 = \beta$, and

$$\beta_2 = \begin{cases} \delta\beta & \text{if } d_1 = 1 \\ \beta & \text{if } d_1 = 0, \end{cases}$$

with $\delta > 0$. This captures the idea that first-period dissent affects the faction's future resource share within the party. For example, if $\delta \in (0, 1)$, dissenting depletes the future resources of the faction (e.g., because of retaliation by the leadership).

The timing of the game is as follows:

1. First period

- The incumbent observes ω_1 and chooses x_1 .
- The faction observes ω_1 and x_1 and chooses d_1 .
- The voter observes $\tilde{x}_1$, updates on γ , and votes.

2. Second period

- The incumbent observes ω_2 and chooses x_2 .
- The faction observes ω_2 and x_2 , and chooses d_2 .
- The policy is implemented.

To guarantee that, for δ sufficiently high, in equilibrium both faction and leader are always interested in being re-elected (regardless of the type of faction), we assume the following:

Assumption 1. $(1 - \beta)R > \pi f(\phi_L, 1)$.

Assumption 2. $\beta R > \pi g(\phi_H, 1)$.

We solve for Perfect Bayesian Nash Equilibria in pure strategies. The incumbent's strategy is a mapping $\sigma^j : \{\phi_j \times \omega_t\} \rightarrow x_t$, such that:

$$\sigma_{\phi_j, \omega_t} = \Pr(x_t = \omega_t | \phi_j, \omega_t).$$

The faction's strategy is a mapping $d_t : \{\phi_I \times \omega_t \times x_t\} \rightarrow d_t$. The voter's strategy is characterized by a re-election rule $\rho(\tilde{x})$, which specifies the probability of re-electing the incumbent after observing the implemented policy $\tilde{x}_1$.

3.1. Discussion of Model

We deliberately align the voter’s ideological preferences with party leadership rather than factions, effectively positioning factions as detrimental to voter interests when reforms are necessary. This assumption creates a more challenging environment for our central finding that increased expected factionalization can improve welfare. However, we do not need perfect alignment between voter and leader. The critical mechanism requires only that voters retain incumbents who demonstrate strength and remove those exhibiting weakness. Our findings remain generalizable even when accounting for voter ideological biases, provided this fundamental electoral behavior persists.

Beside a divergence in preferences for reforms, there could be other reasons for factions to dissent. For instance, dissenting factions might gain opportunities to supplant the current leadership. Under such circumstances, dissent becomes more attractive as ideological incentives are reinforced by the potential gains from leadership succession. This mechanism would expand the range of parameters under which factions choose to dissent in equilibrium, similar to the effect of increasing δ (protecting factions’ resource shares from retaliation) but operating through enhanced future benefits rather than reduced costs of dissent.

In our setting, the faction and the leadership are mis-aligned in terms of policy preferences. This is to capture one of the fundamental reasons for factions to exist: that is, to allow for diverging preferences while sharing the same re-election motives. However, our results do not rely on this assumption. We could relax it to allow for a continuous level of alignment between faction and leader’s preferences. In this more general setting, when the leader and the faction’s preferences are completely aligned, dissent does not happen. However, as long as the faction disagrees sufficiently with the leadership, it will always dissent in equilibrium (provided δ is sufficiently high), despite sharing the same re-election motives. Interestingly, there can be dissent even if the faction agrees with the direction of the reform, but disagrees about the extent to which the reform should go.

Finally, we assume that the effect of factional dissent “scales down” the policy outcome, compared with the proposal. This is a parsimonious way to capture the idea that, if there is dissent, the final outcome of a reform is affected by an intra-party bargaining process, whose result depends on the bargaining weight of different players (therefore, different ϕ_I reflect the different relative bargaining power of the leadership, vis-a-vis factions). However, the main

results would be qualitatively unchanged if, instead, we were to assume that factional dissent reduces the probability that the proposed policy is implemented, and this probabilistic reduction is stronger for weaker leaders. Even in this case, we can find conditions for an over-reform equilibrium, where factions are active and weak leaders behave better than strong leaders in the first period. The non-monotonicity in the welfare function of the voter arises in a similar way.

4. Analysis

We begin with the second period analysis. Our first result shows that, in $t = 2$, the faction always dissents in equilibrium and the elected incumbent party leader $j \in \{I, C\}$ chooses the policy that matches the state of the world. Note that this is true for every δ .

Lemma 1. *In every equilibrium, in the second period:*

- *the faction always dissent;*
- *the incumbent leadership always chooses $x_2 = \omega_2$.*

In the second period there are no electoral concerns. Thus, the faction always dissents, irrespective of ω_2 , as this brings the implemented policy closer to its preferred outcome, increasing its payoffs. The leadership anticipates this, but it remains optimal to choose the policy that matches the state of the world, as this minimizes their policy loss even after factional obstruction.

Given this second-period behavior, we have that the second-period expected payoff of the voter is:

$$\mathbb{E}_{\omega_2, \phi_j}(u_2^v) = -\pi \mathbb{E}_{\phi_j}[f(\phi_j, 1)].$$

Since voters value effective policy implementation, they prefer stronger leaders, who can better overcome factional obstruction. Given this, in the first period V re-elects I if and only if:

$$\Pr(\phi_I = \phi_H | \tilde{x}_1) \geq \gamma.$$

The implemented policy serves as a signal of the leader's type: note that, trivially, $\rho(\tilde{x}_1 = \phi_H) = 1$ and $\rho(\tilde{x}_1 = \phi_L) = 0$.⁶

The first-period analysis is more complex because both leaders and factions must balance immediate policy considerations against future electoral and resource allocation consequences.

⁶In Appendix B, we show that our main results are robust to considering probabilistic revelation.

Below, we organize the analysis of the first period as follows: first, we consider as a benchmark the case where the share of office rents obtained by the faction is independent of the decision of dissenting (i.e., $\delta = 1$). In this scenario, the intra-party competition channel is shut down and dissent primarily operates through the signalling and policy channels. We refer to this benchmark as the *external competition* case.

Then, we unpack the *internal* competition mechanism and consider how the dissent decision by the faction changes when dissenting impacts the faction's share of future rents (i.e., $\delta < 1$), possibly depleting it. This introduces a crucial trade-off for factions: they must weigh the immediate policy benefits of dissent against the potential loss of future resources within the party.

4.1. Benchmark: External Competition

We now analyze the first period behavior. We begin by setting $\delta = 1$, so that dissent does not deplete the faction's future resource share. Hence, we restrict the analysis to the case where dissent does not affect the intra-party balance of power (between leader and faction).

Lemma 2. *In every equilibrium, we have that $d_1 = 1 \forall \phi_I, \omega_1$.*

The intuition for universal dissent is straightforward once we consider each faction type's incentives. Consider a weak faction (facing a strong leader). In this case there is no trade-off: by dissenting the factions pays a lower policy cost, moving the implemented policy $\tilde{x}_1 = \phi_H x_1$ closer to their preferred outcome, and at the same time improves the re-election chances of the party, since the voter re-elects upon observing ϕ_H . Dissent thus serves a dual purpose: policy moderation and credible signaling of leadership strength.

Consider next a strong faction (facing a weak leader). Here, the strategic logic shifts but the conclusion remains the same. Because of the equilibrium behavior of the weak faction, if the voter were to observe $\tilde{x}_1 = 1$ (i.e., the outcome of no dissent), she would negatively update about the strength of the incumbent leader—inferring that only a weak leader could implement policy without obstruction—and therefore would not re-elect. In this case, the faction would be ousted no matter what, and therefore dissents to minimize the policy loss. The next result ensures that an equilibrium where the faction always dissents exists as long as off-path beliefs satisfy D1 (Fudenberg and Tirole, 1991).⁷

⁷Since observing $\tilde{x}_1 = 1$ (no dissent) is off-equilibrium, any belief is consistent with Bayes' rule. However, D1 requires that voters attribute this deviation to the type most likely to benefit from it. Since only strong factions

Lemma 3. *An equilibrium with $d_1 = 1 \forall \phi_I, \omega_1$ always exists, as long as off-path beliefs satisfy D1.*

Having established that factions always dissent in every equilibrium in this benchmark, we now turn to the incentives of the leader to fully characterize the equilibrium. Since dissent is guaranteed, the leader's choice of x_1 directly determines the signal $\tilde{x}_1 = \phi_I x_1$ that voters observe. To ease the notation, in what follows we refer to a strong incumbent party's leader (or simply "incumbent") strategy as σ_{H, ω_1} and a weak incumbent party's leader strategy as σ_{L, ω_1} . In the first period, the following holds:

Lemma 4. *In every equilibrium, $\sigma_{H,1} = 1$ and $\sigma_{L,0} = 1$.*

Intuitively, by implementing the reform when it is optimal to do so ($\omega_1 = 1$), the incumbent fully reveals its type and matches the state. Thus, a strong incumbent faces no trade-off in this case. Similarly, a weak incumbent is always able to match the state of the world without revealing its weakness when $\omega_1 = 0$.

The two potentially ambiguous cases are (i) when a reform needs to be implemented and the incumbent leadership is weak, and (ii) when the status quo needs to be preserved and the incumbent leadership is strong. In both cases, leaders face a tension between matching the state and managing their electoral prospects. The next result shows that a weak incumbent leadership ($\phi_I = \phi_L$) always matches the state of the world.

Lemma 5. *In every equilibrium, $\sigma_{L,1} = 1$.*

This result follows from the fact that, in equilibrium, the voter only re-elects the incumbent with positive probability upon observing $\tilde{x}_1 = 0$ if both types of leaders pool on always matching the state. If instead different types separate, Lemma 4 implies that the strong incumbent leadership never chooses policy 0 when the state is 1, and the weak leader never chooses policy 1 when the state is 0. Therefore, it must be that $\tilde{x}_1 = 0$ is more likely to result from the policy implemented by a weak incumbent. As a consequence, in every non-pooling equilibrium the voter ousts the incumbent upon observing $\tilde{x}_1 = 0$, and the weak leader matches the state $\omega_1 = 1$ since,

(facing weak leaders) might potentially gain from not dissenting, voters must believe that $\tilde{x}_1 = 1$ signals a weak leader. This belief, in turn, eliminates any incentive for strong factions to deviate, as it would guarantee electoral defeat. Thus, D1 ensures the stability of the "always dissent" equilibrium by making any deviation self-defeating for all possible parameters.

conditional on that state, it has no way of getting re-elected. Without any electoral benefit from deviation, the weak leader defaults to the policy-optimal choice.

We can summarize this observation as follows:

Remark 1. *In equilibrium, $\rho(0) > 0$ if and only if $\sigma_{I,\omega} = 1$ for all ϕ_I and ω .*

This remark captures a key insight: voter skepticism toward low reform outcomes can only be overcome if both leader types always match the state. Given the results above, we now focus our analysis on the ϕ_H type's trade-off when $\omega_1 = 0$: on the one hand, implementing a reform reveals that the incumbent leadership is strong, therefore it secures re-election. On the other hand, the reform is not optimal given the state. This is where the perverse incentive for over-reform emerges: strong leaders may sacrifice policy optimality for electoral gain.

4.2. Equilibria

We begin by asking whether it is possible to sustain an equilibrium with “full discipline,” in which parties in power always choose the optimal policy, and the voter re-elects the incumbent with any probability.⁸ Note that it follows from Lemma 2 that factions always dissent in every equilibrium. To simplify the notation in the results that follow, we omit the strategy of the faction and the voter re-election rule upon observing $\tilde{x}_1 = 1$.

Proposition 1. *There exists a full discipline equilibrium where*

- (i) $\sigma_{I,\omega} = 1$ for every ω, ϕ_I , and
- (ii) $\rho(0) = 1$ if and only if $(1 - \beta)R \leq f(0, 1) - f(\phi_L, 1) + \pi f(\phi_L, 1)$;
- (iii) $\rho(0) = 0$ if and only if $(1 - \beta)R \leq f(\phi_H, 0) + \pi f(\phi_H, 1)$.

There can be two types of full discipline (pure strategy) equilibria, depending on the retention rule of the voter. They have different incentive compatibility conditions. Proposition 1 (ii) derives the incentive compatibility condition for a full discipline equilibrium with $\rho(0) = 1$. In this case, it is the weak incumbent leadership that may have a profitable deviation from full discipline, as the strong incumbent leadership is always re-elected. The weak incumbent leadership, instead, is voted out if $x_1 = 1$. Therefore, in this equilibrium, when being re-elected is very valuable, a weak incumbent leadership has incentive to deviate from a full discipline equilibrium when $\omega_1 = 1$: while it pays a cost from mismatch today, it gets re-elected and enjoys R tomorrow. As long as

⁸Notice that in this equilibrium the voter is indifferent upon observing 0, since it is a pooling equilibrium. Therefore, while we focus on pure strategy equilibria, there are also mixed strategy equilibria with $\rho(0) \in (0, 1)$.

rents from office are not too high, this condition is satisfied. Similarly, Proposition 1(iii) derives the relevant incentive compatibility condition for a full discipline equilibrium with $\rho(0) = 0$. In this case, it is the strong incumbent leadership that may have the incentive to deviate. When $\omega_1 = 0$, the incumbent can implement a reform that reveals its strength and ensures re-election, while in equilibrium it loses re-election by choosing $x_1 = 0$. As long as rents are lower than the policy cost from mismatch, the strong leader behaves well, matching the state in equilibrium.

It follows from Proposition 1 that for some parameter values we cannot sustain full discipline in equilibrium. In other words, there must be cases in which the incumbent party decides not to adopt the right reform. Because of Lemma 4 and Lemma 5, we know that if such incentive exists, it comes from the strong leader, who decides to implement a reform when not needed (i.e., when $\omega_1 = 0$). We refer to this behavior as “over-reforming”. The next result shows under what conditions the equilibrium features over-reforms.

Proposition 2. *There exists an “over-reform” equilibrium where*

- (i) $\sigma_{H,0} = 0$ and $\sigma_{I,\omega} = 1$ otherwise;
- (ii) $\rho(0) = 0$ if and only if $(1 - \beta)R \geq f(\phi_H, 0) + \pi f(\phi_H, 1)$.

The equilibrium is unique for sufficiently high values of R .

When the value of office is sufficiently high, strong leaders—those facing weak factions—implement unnecessary reforms to signal their control over the party. This over-reform strategy secures re-election but comes at the cost of policy optimality. The key insight of Proposition 2 is that strong factions serve as a disciplining device, even though they are misaligned with voter preferences. Leaders facing strong factions (weak leaders) cannot get re-elected by reforming because factional obstruction severely limits policy implementation. Knowing they will be ousted regardless of their actions—voters correctly infer that $\tilde{x}_1 = 0$ likely indicates a weak leader—these leaders default to choosing the optimal policy. In contrast, leaders facing weak factions (strong leaders) can effectively signal their type by implementing reforms that pass with minimal obstruction. This creates a perverse incentive: precisely because they *can* overcome factional resistance, they *do* implement reforms even when maintaining the status quo would be optimal. Thus, strong factionalization paradoxically improves policymaking by removing the temptation to over-reform: weak leaders behave well not despite their weakness, but because of it.

Figure 1 plots the different equilibria as a function of R and ϕ_H . Notice that the orange region identifies the parameter space where the over-reform equilibrium is unique, that is for values of office rents sufficiently high.

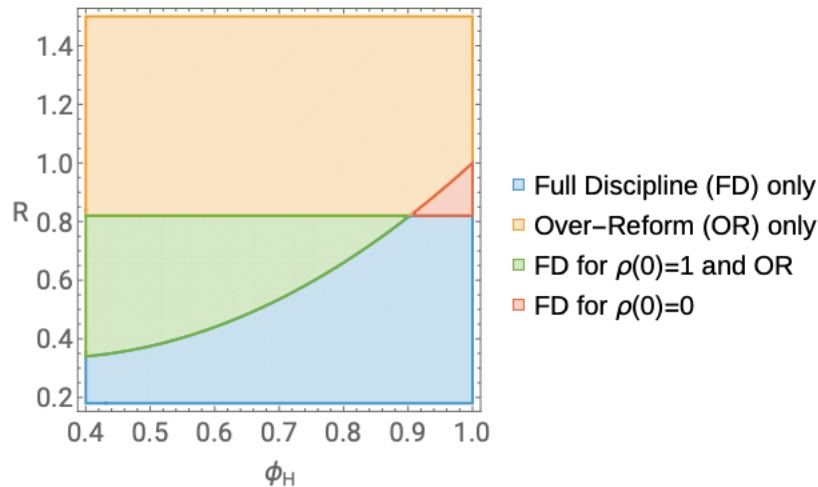


Figure 1 – Different equilibria for ϕ_H and R , assuming $f(\tilde{x}, \omega) = (\tilde{x} - \omega)^2$. Other parameters are set to $\phi_L = 0.4$, $\beta = 0.1$ and $\pi = 0.5$. As displayed in the legend, for any retention rule adopted by the voter in equilibrium, the blue region on the bottom of the panel corresponds to an area with multiple Full-Discipline equilibria (both with $\rho(0) = 0$ and $\rho(0) = 1$), whereas the orange region on the top corresponds to the Over-Reform equilibrium only. For intermediate values of R , different $\rho(0)$ can be associated with different behaviors by the Incumbent in equilibrium. On the green region there exists both a Full-Discipline equilibrium (with $\rho(0) = 1$), and an Over-Reform equilibrium (with $\rho(0) = 0$). The red region corresponds to a Full-Discipline equilibrium when $\rho(0) = 0$.

4.3. Discussion of Results

We conclude this section with a discussion of the robustness of the over-reform equilibrium and the strategic role of factions. Appendix B formally demonstrates that the equilibrium does not depend on the assumption that $x_t = 1$ perfectly reveals the incumbent leader's type. Revelation can be probabilistic, and it is always possible to choose a sufficiently high (strictly less than one) probability of revelation such that the equilibrium continues to hold.

Moreover, the equilibrium does not rely on the voter's inability to observe $\tilde{x}_1$ or u_1^v (i.e., the consequences of a particular policy) before the election. The voter always faces a selection problem going into the second period. This implies that even a poorly performing reform can enhance the incumbent party's re-election prospects, as long as it signals that $\phi_I = \phi_H$. What

matters is not the policy’s actual consequences but what its successful implementation reveals about the leader’s ability to overcome factional resistance. Additionally, since observing $\tilde{x}_1$ allows the voter to perfectly infer x_1 , given the knowledge of factional dissent, any assumptions about the direct observability of x_1 are irrelevant to our results. Finally, Appendix C shows that the over-reform equilibrium and the core logic behind our welfare findings remain robust even when we allow the prior about the leadership’s strength (γ) to differ between the challenger and the incumbent.

The central mechanism driving our equilibria is that factional obstruction creates an informational channel whereby implementing a reform reveals information about the leadership’s type, which in turn influences voter behavior in re-election decisions. Voters reward signals of strong leadership—demonstrated by the ability to push reforms through factional resistance—even if the policy implemented is suboptimal.

This mechanism is consistent with empirical evidence showing increased support for Trump in regions economically harmed by the trade war (Autor et al., 2024). This is also consistent with the empirical results of Cipullo and Lee (2025): districts hit by the “China shock” reward electorally more effective legislators. Conversely, when reforms reveal leadership weakness or party factionalization, voters respond negatively. A notable example is Theresa May’s proposed “Brexit deal,” which was defeated in Parliament due to widespread defections within her own Tory party. Her approval rating dropped sharply from -30% in January 2019—before the parliamentary votes—to -45% by March 2019, after the votes (The Guardian, 2019). The visible factional revolt signaled May’s inability to control her party, triggering electoral punishment.

Finally, this first set of results reveal three key insights about factional activism. First, when dissent does not affect the balance of power within parties, factional dissent is a pervasive phenomenon that emerges irrespective of the strength of factions—even weak factions dissent because it improves both their policy payoffs and their party’s electoral prospects. Second, factions may have incentives to act even if this implies an electoral loss for the party. This is consistent with equilibrium behavior, since strong factions (facing weak leaders) know their party cannot secure re-election regardless of their actions, so they prioritize immediate policy gains. Third, this universal dissent creates the informational environment that enables over-reform: because factions always obstruct, the degree of policy implementation perfectly reveals leadership strength.

4.4. Voter Welfare

In this section we study the welfare effect of expected factionalization and of over-reforms in equilibria where factions are always active. We first note that, under a full-discipline equilibrium, voter welfare is trivially increasing in both γ and ϕ_H . Since both types of incumbent leadership always implement the correct reform, the more likely it is that the incumbent leadership is strong (and the more likely it is able to implement the chosen policy), the better it is for the voter.

We now study the welfare effects of expected party factionalization in the over-reform equilibrium. Crucial to our model is the voter uncertainty about how divided the incumbent party is. It is this uncertainty that drives the key strategic forces in the over-reform equilibrium. Therefore, it is natural to measure expected factionalization with γ , i.e., the prior probability of the incumbent leadership being strong (low γ implies high factionalization). The next result shows that voter welfare can either be increasing or decreasing in γ .

Proposition 3. *Voter welfare (W) can be increasing or decreasing in factionalization (γ). Furthermore, W is non-monotonic in γ if and only if:*

$$\frac{(1 - \pi)f(\phi_H, 0)}{\pi(f(\phi_L, 1) - f(\phi_H, 1))} \in (1, 3).$$

Proposition 3 shows that strong, active factions can be good precisely because they are bad. Why? The voter wants to punish factionalization, and she does so given the retention rule used in equilibrium, which prescribes to elect the challenger when the observed policy is $\tilde{x}_1 = 0$. Given this, the incumbent leadership of a high-factionalized incumbent party ($\phi_I = \phi_L$) can only behave well in equilibrium, always matching the state. This is because it will be ousted no matter what, and at least it can choose to implement the correct policy. In contrast, the incumbent leadership of a low-factionalized incumbent party ($\phi_I = \phi_H$) faces a trade-off when $\omega_1 = 0$. On the one hand, it wants to match the state of the world. On the other, it wants to signal it is not factionalized by choosing $x_1 = 1$. In this equilibrium, the latter effect prevails and therefore the low-factionalized party implements reforms even when not needed.⁹

Thus, in this equilibrium γ has competing effects on the voter welfare. On the one hand, a low-factionalized party is better in implementing the reform, when it is necessary to do so

⁹We note that the trade-off on voter welfare that emerges in this model is different from the usual trade-off between first period behavior and selection (see, e.g., [Besley \(2006\)](#); [Ashworth and Bueno De Mesquita \(2014\)](#); [Trombetta \(2020\)](#)). In our case, selection is perfect and the trade-off is between “efficiency” and good behavior.

(i.e., $f(\phi_H, 1) < f(\phi_L, 1)$). On the other hand, a low-factionalized party has the incentive to over-reform, therefore it is less likely to match the state of the world when $\omega_1 = 0$. These two forces push the welfare effect of expected factionalization in opposite directions, and it can be that the welfare is non-monotonic, i.e. reverse-u shaped, in γ .¹⁰ Formally, the derivative of W with respect to γ is negative (i.e., higher factionalization is good for the voter) when the following holds:

$$\underbrace{(3 - 2\gamma)}_{\text{Benefit from selection}} \underbrace{\pi [f(\phi_L, 1) - f(\phi_H, 1)]}_{\text{Benefit from better reform}} < (1 - \pi)f(\phi_H, 0). \quad (1)$$

To understand the intuition, note that while the cost of over-reform (i.e., $(1 - \pi)f(\phi_H, 0)$) is constant in γ , the welfare benefit of having a low-factionalized party is decreasing in γ (the LHS in (1)). We can separate this benefit in two components: one due to the better implementation of a reform when the reform is needed, and one due to better selection (note that the strong incumbent leadership is always re-elected, in the separating equilibrium). The benefit from selection is decreasing in γ , because a high γ implies that it is more likely that a strong faction is replaced by another (challenger) strong type. Hence, when γ is low, the selection effect may push the benefit from a strong incumbent leadership above its cost (the mismatch when $\omega_1 = 0$). When γ increases, however, the cost effect may start dominating. Finally, equation (1) highlights the role of π . The higher is the probability that the reform is needed, the less likely is condition (1) to hold (i.e., it is less likely that strong factions are good for the voter). This is intuitive: strong leaders are better precisely when they implement a needed reform. Moreover, an increase in π also decreases the probability of observing an over-reform in period 1.

The second part of Proposition 3 shows the conditions that guarantee a non-monotonicity in voter welfare. Formally, the ratio between the policy cost and the policy benefit of having a strong incumbent leadership, weighted by the probability of the state where those costs or benefits materialize, should be intermediate. Bigger than 1, because the welfare cost of the strong incumbent leadership is “paid” only in one period, while the benefit applies potentially to both periods, but it cannot be too large (otherwise the welfare is always decreasing). Finally, it is important to notice that the optimality of an interior γ relies on γ being strictly smaller than 1. If $\gamma = 1$ there is no over-reform and the voter achieves the first best outcome.

¹⁰It is easy to show that voter welfare can also be non-monotonic in ϕ_H when $f(\tilde{x}_t, \omega_t)$ is quadratic, although this result rests on the specific functional form adopted.

While over-reform equilibria exist in the literature, the strategic mechanism driving our non-monotonic welfare result represents a novel contribution. Our model generates a distinctive inversion of the standard selection-control trade-off. In typical political agency models, improved selection (electing high-quality politicians) comes at the cost of worse behavior by low-quality types, who engage in inefficient signaling to mimic high types. Our model exhibits the opposite pattern: electing high-quality leaders (those with strong party control) comes at the expense of bad behavior by the high types themselves. This inversion occurs because strong leaders systematically behave worse than weak leaders in our setting—implementing unnecessary reforms to signal their strength, while weak leaders, knowing they cannot credibly signal, default to optimal policy choices. It is precisely this counter-intuitive result—that the most capable politicians engage in the most distortionary behavior—that drives the non-monotonicity in voter welfare with respect to expected party factionalization.

So far we focused on voter welfare within the over-reform equilibrium. An interesting question is whether this equilibrium, characterized by some policy distortion, can be better, for the voter, than the full discipline equilibria. The reader may think at this point that the latter are always better, because parties always choose the optimal policy. In fact, we show that this is not always the case.

Proposition 4. *Voter welfare is higher in the over-reform equilibrium than in the full discipline equilibrium if $\pi [f(\phi_L, 1) - f(\phi_H, 1)] > f(\phi_H, 0)$ and γ is sufficiently low.*

The comparison between the two equilibria boils down to the difference in expected payoffs when $\omega_1 = 0$ and $\phi_I = \phi_H$. In the over-reform equilibrium, the strong incumbent leadership implements a damaging reform today and is re-elected, guaranteeing a higher payoff tomorrow.

Welfare in the full discipline equilibrium is a function of the voter retention rule. When $\rho(0) = 0$, the strong incumbent leadership chooses the correct policy (i.e. no reform when $\omega_1 = 0$) but is removed from office. If the replacement is expected to be strong with sufficiently high probability, the latter is always better for the voter. But, when γ is sufficiently low, and the gains from a correct policy when $\omega_2 = 1$ are sufficiently high, the over-reform equilibrium dominates the full discipline one.

When $\rho(0) = 1$, the strong incumbent leadership always remains in office (as in the over-reform equilibrium). Therefore, the difference in welfare under the two equilibria depends on the policy outcome produced when $\omega_1 = 0$ (unnecessary reform in the over-reform equilibrium

and no reform in the full discipline equilibrium). Moreover, the weak incumbent leadership is retained when $\tilde{x}_1 = 0$ in the full discipline equilibrium, while it is replaced in the over-reform equilibrium. Once again, the overall solution of this comparison is not trivial, but the over-reform equilibrium is more likely to be superior for the voter when γ is small, because the cost of retaining a highly-factionalized party would be high in a full discipline equilibrium.

4.5. When do ‘Good’ Politicians Benefit Voters?

Our analysis reveals a counter-intuitive finding: strong, active factions—often view as obstacles to effective governance—can actually enhance voter welfare. While we begin by assuming that factional constraints impede policy implementation, our model demonstrates that this very impediment serves as a commitment device to exert policy-making restraint. When party leaders’ interests diverge from voters’ preferences, factions are obviously beneficial. More surprisingly, we show that stronger factions can enhance welfare even when leaders are aligned with voters, by preventing less capable leaders from implementing excessive reforms.

The welfare benefits of factionalization, in a context where factions have reasons to be always active, depend critically on the distribution of party types in the political landscape. When strong, unified parties are relatively rare (low γ), voters are willing to tolerate unnecessary reforms in exchange for effectiveness when reforms are genuinely needed. However, as strong, unified parties become more common (high γ), the voter’s tolerance for such signalling-motivated over-reforms diminishes. This generates our distinctive non-monotonic welfare result: moderate levels of expected factionalization may optimize voter welfare by balancing the trade-off between effective implementation and disciplined policy selection.

This non-monotonicity complements existing findings in the literature on pandering. In particular, [Kartik and Van Weelden \(2019\)](#) show that welfare may be minimized at intermediate prior beliefs (on the incumbent being a good type), as these create the strongest pandering incentives for good-type politicians seeking to distinguish themselves from biased types. Our model yields the opposite pattern—a reverse-U shaped welfare function—because of a fundamentally different mechanism. In [Kartik and Van Weelden \(2019\)](#)’s framework, the intensity of pandering incentives varies with the prior probability on politician types, creating a welfare depression at intermediate beliefs. When politicians are believed to be either mostly good or mostly bad, pandering incentives weaken—in the former case because reputation concerns are less pressing, and in the latter because reputation-building becomes too costly given voters’ skepticism.

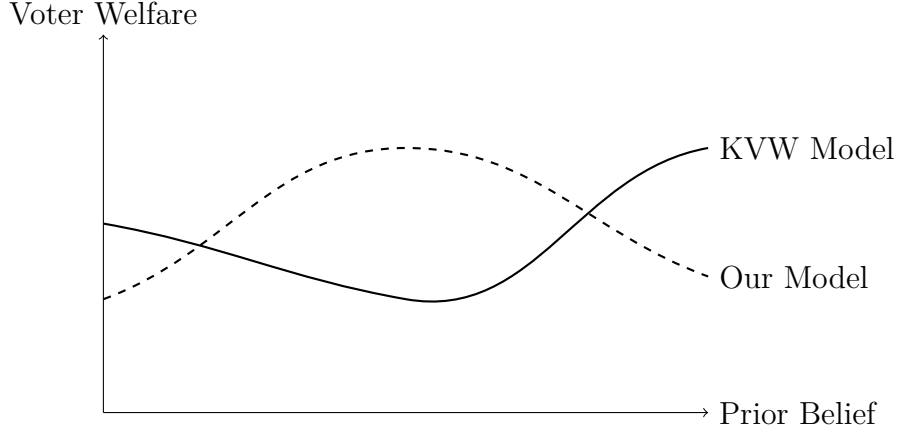


Figure 2 – Comparative Welfare Results: Our model produces a reverse U-shaped welfare curve where moderate levels of leadership strength (intermediate values of γ) maximize voter welfare. In contrast, [Kartik and Van Weelden \(2019\)](#)’s model yields a U-shaped welfare curve where intermediate priors (probability p that politicians are congruent) minimize welfare by creating the strongest pandering incentives.

By contrast, in our model, politicians’ incentives to choose the wrong policy remain constant across priors, but the *social value* of this behavior varies with the prior. This generates the reverse-U shaped welfare function shown in Figure 2. When strong, unified parties are rare (low γ), voters benefit from the informational value of identifying them, making the occasional unnecessary reform an acceptable cost. When strong parties are common (high γ), voters face less uncertainty about implementation capability, so the welfare cost of over-reform dominates. At intermediate levels of party factionalization, voters experience an optimal balance: sufficient certainty about policy implementation without excessive signaling-driven reforms.

This difference highlights a novel insight about institutional design: while welfare in [Kartik and Van Weelden \(2019\)](#)’s model is maximized at extreme prior beliefs (where pandering pressures are muted), welfare in our model peaks at moderate levels of leadership strength (where implementation effectiveness is balanced with restrained reform impulses).

4.6. Intra-Party Competition

We now explore how factional dissent is affected when intra-party competition introduces dynamic considerations for factions. Specifically, we relax the assumption that $\delta = 1$, allowing a faction’s dissent to impact its future share of party resources. When δ is sufficiently close to 1, factions continue to dissent as in our benchmark model. In the second period, irrespective of δ ,

both leaders and factions maintain the same behavior pattern regardless of δ , as second-period decisions involve no further strategic considerations about future payoffs.

Our analysis now focuses on identifying conditions under which factions might strategically refrain from dissent ($d_1 = 0$). We first establish an important asymmetry between faction types:

Lemma 6. *There is no equilibrium such that $d_1 = 1$ for $\phi_I = \phi_H$ and $d_1 = 0$ for $\phi_I = \phi_L$.*

This result reveals the distinct strategic incentives facing different faction types. If we conjecture an equilibrium where only weak factions dissent, voters would rationally never re-elect upon observing no dissent (i.e., $\rho(1) = 0$), correctly inferring this signals a strong faction. Since strong factions' parties cannot secure re-election through non-dissent, they have no incentive to refrain from dissent—which would at least improve their policy payoff.

When dissent reduces a faction's future rents ($\delta < 1$), we can identify equilibria where factions strategically withhold opposition despite policy disagreements:

Proposition 5. *Let $\delta \in [0, 1)$. There exists an over-reform equilibrium where:*

- (i) $\sigma_{H,0} = 0$ and $\sigma_{I,\omega} = 1$ otherwise;
 - (ii) $\rho(0) = 0$ and $\rho(1) = 1$;
 - (iii) no faction dissents;
- for appropriate parameter values.

Why do factions remain silent in this over-reform equilibrium? The key intuition is that when $\delta < 1$, dissent carries a future cost—factions that oppose the leader receive fewer resources in subsequent periods. In the over-reform equilibrium, voters re-elect parties that appear united (no observed dissent), interpreting unity as a signal of party strength. This creates a coordination problem: individual factions prefer to dissent on policy grounds, but collectively benefit from appearing united to secure re-election. As long as the benefits from re-elections are sufficiently high, factions find it optimal to remain silent despite policy disagreements, prioritizing their long-term rents over immediate policy preferences.

The strategic calculus governing factional dissent fundamentally depends on how party resources are allocated. When party leaders maintain significant discretion over factional rewards, dissent carries greater risk as factions' future prospects become contingent on their present loyalty. Conversely, when rewards are determined by more objective measures insulated from leadership discretion—such as vote shares for geographically concentrated factions or membership

contributions— factions can more freely express disagreement without fear of retribution and therefore they are always active, as shown in the equilibria discussed in Section 4.1.

Empirical Implication 1. *We should expect more active factional dissent in more institutionalized parties, where intra-party allocation mechanisms are transparent and protected from leadership manipulation.*

We now compare voter welfare across two variants of the over-reform equilibrium: one where factions actively dissent (as in our baseline model with $\delta = 1$) and another where factions remain strategically silent (which can happen when $\delta < 1$). While one might expect that the absence of factional dissent would benefit voters (given that factions in our model are both misaligned with voter preferences and reduce policy effectiveness), our analysis reveals a more nuanced reality. Under specific conditions, factional dissent can actually enhance voter welfare, even when comparing the same equilibrium type (over-reform) with and without active factions.

Proposition 6. *There exists a set of parameters such that welfare in the over-reform equilibrium is higher when factions are active than when factions are inactive.*

This cross-equilibrium comparison between over-reform scenarios—with and without factional activism—reveals three competing effects. First, factional activism creates a cost in terms of “policy attenuation:” when a needed reform is proposed, the final outcome is better for the voters if factions do not interfere with the process. However, factions are also beneficial for the voter for two reasons. First, there is a selection motive: when factions are active, the voter is better able to spot weak leaders and replace them. Second, when factions are active, over-reform equilibria are less damaging because their effectiveness is reduced by factional activism. Therefore, dissent acts as a beneficial constraint when reforms are seldom necessary. Since strong leaders implement unneeded reforms to signal their strength, factional dissent effectively constrains this welfare-reducing behavior. When reforms are rarely optimal (π is low) and the risk of over-reform is high (γ is large), this constraint provides substantial benefits. Figure 3 indeed shows that, when γ is sufficiently high (strong leaders are common) and π is sufficiently low (reforms are rarely needed), an equilibrium with factional dissent improves voter welfare compared to one without dissent, both conditional on leaders following an over-reform strategy.

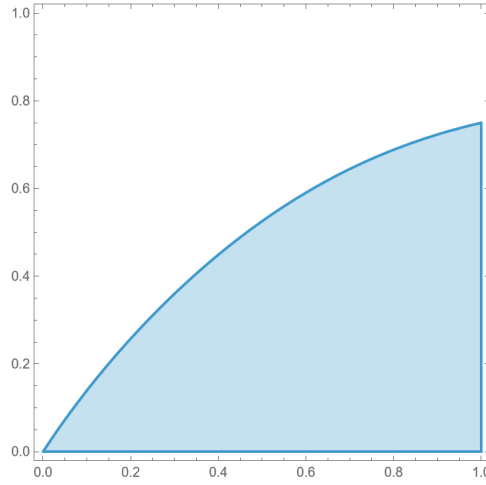


Figure 3 – The blue region highlights the parameter space such that the Over-Reform equilibrium with active factions is welfare superior with respect to the Over-Reform equilibrium without active factions, as a function of γ (x axis) and π (y axis). We assume $f(\tilde{x}, \omega) = (\tilde{x}_t - \omega_t)^2$. Other parameters are set to $\phi_L = 0.3$, $\phi_H = 0.5$.

5. Conclusion

In this paper we develop a theory of policy-making in the presence of factions. Contrary to the conventional wisdom that views factions primarily as impediments to effective governance, our analysis reveals them to be strategic actors whose dissent shapes both policy outcomes and electoral dynamics, demonstrating that factions serve a crucial disciplining role in the political process, even when they are misaligned with voter preferences. Our results highlight the strategic reasons behind factional activism (even when this implies a reputational cost for the party), and also the conditions under which they improve voter’s welfare.

Our key finding is that factions create an informational environment that fundamentally alters leadership incentives. When factions actively dissent—obstructing policy implementation proportionally to their strength—they enable voters to infer leadership quality from observed policy outcomes. This informational channel generates a perverse result: strong party leaders may implement unnecessary reforms to signal their control over the party. This “over-reform” equilibrium emerges because voters, observing minimal factional obstruction, interpret successful reform implementation as evidence of leadership strength worth rewarding electorally.

In contrast, strong factions prevent over-reform. Leaders facing powerful factional opposition cannot credibly signal their strength through unnecessary reforms because factional obstruction would severely dilute any policy they attempt. Knowing they face electoral punishment regard-

less, these leaders default to optimal policy choices. Thus, factions improve policy selection not through deliberation or information aggregation, but through strategic constraint.

This mechanism produces the counterintuitive result that voter welfare can be non-monotonic in the expected level of party factionalization. While weak factions enable better policy implementation when reforms are genuinely needed, they also fail to prevent unnecessary reforms undertaken for signaling purposes. Consequently, an intermediate level of expected factionalization optimizes welfare by balancing two opposing forces: sufficient factional strength to constrain excessive reforms while maintaining adequate leadership capacity to implement necessary policy changes.

Our analysis also reveals when factions choose to exercise their constraining role. We show that factional dissent depends critically on intra-party institutions. When factions' resource shares are protected from leadership retaliation, they dissent actively, creating the disciplining effect we identify. However, when dissent jeopardizes future rents, factions may strategically acquiesce, enabling both beneficial reforms and harmful over-reform. This suggests that the design of intra-party allocation mechanisms—whether resources follow objective rules or leadership discretion—fundamentally shapes factions' ability to constrain excessive policy-making.

These findings reframe our understanding of party organization in democracies. Rather than viewing factions as obstacles to be overcome, our theory suggests they are essential features that can enhance democratic accountability. The welfare benefits of factionalization challenge reformers who seek to strengthen party discipline and leadership control. Indeed, our results imply that institutions protecting factional autonomy—such as transparent resource allocation rules and formal faction rights—may improve policy outcomes by enabling factions to serve their constraining function.

References

- Acemoglu, D., G. Egorov, and K. Sonin (2013, May). A Political Theory of Populism*. *The Quarterly Journal of Economics* 128(2), 771–805.
- Ashworth, S. and E. Bueno De Mesquita (2014, August). Is Voter Competence Good for Voters?: Information, Rationality, and Democratic Performance. *American Political Science Review* 108(3), 565–587.
- Ashworth, S. and K. W. Shotts (2010). Does informative media commentary reduce politicians’ incentives to pander? *Journal of Public Economics* 94(11-12), 838–847.
- Autor, D., A. Beck, D. Dorn, and G. H. Hanson (2024). Help for the heartland? the employment and electoral effects of the trump tariffs in the united states. Technical report, National Bureau of Economic Research.
- Besley, T. (2006). *Principled agents?: The political economy of good government*. Oxford University Press.
- Bils, P. and F. Izzo (2025). Policy gambles and valence in elections. *Unpublished Manuscript*. Previously circulated as ”Policymaking in Times of Crisis”.
- Blumenthal, B. (2024, September). Environmental Policymaking with Political Learning.
- Buisseret, P. (2016). “together or apart”? on joint versus separate electoral accountability. *The Journal of Politics* 78(2), 542–556.
- Buisseret, P. and C. Prato (2016). Electoral control and the human capital of politicians. *Games and Economic Behavior* 98, 34–55.
- Caillaud, B. and J. Tirole (1999). Party governance and ideological bias. *European Economic Review* 43(4–6), 779–789.
- Caillaud, B. and J. Tirole (2002). Parties as political intermediaries. *The Quarterly Journal of Economics* 117(4), 1453–1489.
- Canes-Wrone, B., M. C. Herron, and K. W. Shotts (2001). Leadership and pandering: A theory of executive policymaking. *American Journal of Political Science*, 532–550.

- CBS News (2018). Republicans sign letter urging trump to back off on tariffs.
- Che, Y.-K., W. Dessein, and N. Kartik (2013). Pandering to persuade. *American Economic Review* 103(1), 47–79.
- Cipullo, D. and B. E. Lee (2025). Economic shocks and the electoral fortunes of effective legislators. Working Paper 11791, CESifo, Munich, Germany.
- Cox, G. W. and M. D. McCubbins (2007). *Legislative leviathan: Party government in the House*. Cambridge University Press.
- Crutzen, B. S., M. Castanheira, and N. Sahuguet (2010). Party organization and electoral competition. *The Journal of Law, Economics, & Organization* 26(2), 212–242.
- Daley, B. and E. Snowberg (2011, August). Even if it is not Bribery: The Case for Campaign Finance Reform. *Journal of Law, Economics, and Organization* 27(2), 324–349.
- Dessein, W. (2002). Authority and communication in organizations. *Review of Economic Studies* 69(4), 811–838.
- Dewan, T. and R. Hortala-Vallve (2019, July). Electoral Competition, Control and Learning. *British Journal of Political Science* 49(3), 923–939.
- Dewan, T. and F. Squintani (2016, October). In Defense of Factions. *American Journal of Political Science* 60(4), 860–881.
- Fox, J. and M. C. Stephenson (2011). Judicial review as a response to political posturing. *American Political Science Review* 105(2), 397–414.
- Fox, J. and R. Van Weelden (2010). Partisanship and the effectiveness of oversight. *Journal of Public Economics* 94(9-10), 674–687.
- Fox, J. and R. Van Weelden (2012). Costly transparency. *Journal of Public Economics* 96(1-2), 142–150.
- Fu, Q. and M. Li (2014). Reputation-concerned policy makers and institutional status quo bias. *Journal of Public Economics* 110, 15–25.
- Fudenberg, D. and J. Tirole (1991). *Game theory*. MIT press.

- Gratton, G., L. Guiso, C. Michelacci, and M. Morelli (2021, September). From Weber to Kafka: Political Instability and the Overproduction of Laws. *American Economic Review* 111(9), 2964–3003.
- Greene, Z. D. and M. Haber (2015). The consequences of appearing divided: An analysis of party evaluations and vote choice. *Electoral Studies* 37, 15–27.
- Hirsch, A. V. and J. P. Kastellec (2022). A theory of policy sabotage. *Journal of Theoretical Politics* 34(2), 191–218.
- Invernizzi, G. M. (2023). Antagonistic cooperation: Factional competition in the shadow of elections. *American Journal of Political Science* 67(2), 426–439.
- Invernizzi, G. M. and F. Izzo (2025). Evolving parties. *Unpublished Manuscript*.
- Invernizzi, G. M. and C. Prato (2024). Bending the iron law: The distribution of power within political parties. *American Journal of Political Science*.
- Izzo, F. (2024). With friends like these, who needs enemies? *The Journal of Politics* 86(3), 835–849.
- Kartik, N. and R. Van Weelden (2019, March). Informative Cheap Talk in Elections. *The Review of Economic Studies* 86(2), 755–784.
- Lehrer, R., P. Stöckle, and S. Juhl (2024). Assessing the relative influence of party unity on vote choice: evidence from a conjoint experiment. *Political Science Research and Methods* 12(1), 220–228.
- Levy, G. (2004, June). Anti-herding and strategic consultation. *European Economic Review* 48(3), 503–525.
- Lodato, S., C. Mavridis, and F. Vaccari (2024). The unelected hand? bureaucratic influence and electoral accountability. *arXiv preprint arXiv:2402.17526*.
- Madison, J. (1787). The federalist no. 10. *November 22*(1787), 1787–88.
- Majumdar, S. and S. W. Mukand (2004). Policy gambles. *The American Economic Review* 94(4), 1207–1222.

- Merzoni, G. and F. Trombetta (2022, September). Pandering and state-specific costs of mismatch in political agency. *Games and Economic Behavior* 135, 132–143.
- Persico, N., J. C. Pueblita, and D. Silverman (2011). Factions and political competition. *Journal of Political Economy* 119(2), 242–288.
- Prat, A. (2005). The wrong kind of transparency. *American Economic Review* 95(3), 862–877.
- Prato, C. and S. Wolton (2018, October). Electoral Imbalances and Their Consequences. *The Journal of Politics* 80(4), 1168–1182.
- Rohde, D. W. (1991). *Parties and leaders in the postreform House*. University of Chicago Press.
- Shaver, B. (2025). Signaling ability through policy change. Under review (working paper).
- The Guardian (2019). Theresa may: a political obituary in five charts.
- Trombetta, F. (2020). When the light shines too much: Rational inattention and pandering. *Journal of Public Economic Theory* 22(1), 98–145.
- Tsebelis, G. (2002). *Veto players: How political institutions work*. Princeton University Press.
- Wawro, G. J. and E. Schickler (2013). *Filibuster: Obstruction and Lawmaking in the U.S. Senate*. Princeton University Press.

Appendix

Table of Contents

1. Appendix A: Main Results - Proofs (*page 1*)
2. Appendix B: Extension: Probabilistic Type Revelation (*page 10*)
3. Appendix C: Different Degrees of Factionalization (*page 11*)
4. Appendix D: Continuous Dissent Choice (*page 15*)

A. Main Results - Proofs

Proof of Lemma 1. Going backward, consider the decision of any faction after observing $x_2 = 1$. Irrespective of ω_2 , $-g(\phi_I, \omega_2) > -g(1, \omega_2)$, therefore $d_2 = 1$.

Consider the decision of any leader in period 2. As there are no re-election concerns, choosing $x_2 = \omega_2$ maximizes $f(\phi_I, \omega_2) \forall \phi_I$. Therefore, $\sigma_{I, \omega_2} = 1$. $\square$

Proof of Lemma 2. First, recall that in every PBE, $\rho(\phi_H) = 1$ and $\rho(\phi_L) = 0$. Then, we complete the proof in two steps. First, we show that the weak faction (i.e. factions in parties such that $\phi_I = \phi_L$) always dissent. Then, we show that there are no re-election probabilities compatible with sequential rationality such that the strong faction does not dissent.

Claim 1. *In every equilibrium, the weak faction ($\phi_I = \phi_L$) chooses $d_1 = 1$ for every $\rho(\tilde{x}_1 = 1)$.*

Proof of Claim 1. Suppose $\omega_1 = 0$. The (weak) faction's payoff from a deviation to $d_1 = 0$ is:

$$-g(1, 0) + \rho(\tilde{x} = 1)(\beta R - \pi g(\phi_H, 1)), \quad (\text{A.1})$$

whereas the equilibrium payoff from $d_1 = 1$ is:

$$-g(\phi_H, 0) + \beta R - \pi g(\phi_H, 1), \quad (\text{A.2})$$

This follows from the fact that, in every equilibrium, $\rho(\tilde{x} = \phi_H) = 1$. As long as $\beta R > \pi g(\phi_H, 1)$, i.e. as long as the second period payoff is positive and therefore parties have an interest in staying in power, the payoff from $d_1 = 0$ is maximized at $\rho(\tilde{x} = 1)$. Even in that case, (A.2) is larger than (A.1) because $\phi_H < 1$. This implies that in equilibrium $d_{H,0} = 1$, irrespective of the (possibly off path) beliefs that determine $\rho(\tilde{x} = 1) = 1$.

Suppose $\omega_1 = 1$. The faction's payoff in equilibrium from a deviation to $d_1 = 0$ is

$$-g(1, 1) + \rho(\tilde{x} = 1)(\beta R - \pi g(\phi_H, 1)), \quad (\text{A.3})$$

whereas the equilibrium payoff from $d = 1$ is

$$-g(\phi_H, 1) + \beta R - \pi g(\phi_H, 1). \quad (\text{A.4})$$

For the same logic as above, this implies that in equilibrium $d_{H,1} = 1$ for every possible belief following $\tilde{x} = 1$. $\square$

Claim 2. *In every equilibrium, the strong faction ($\phi_I = \phi_L$) chooses $d_1 = 1$.*

Proof of Claim 2. Define E_L as the overall payoff of the strong faction in case of dissent:

$$E_L = \begin{cases} -g(\phi_L, 0) & \text{if } \omega = 0 \\ -g(\phi_L, 1) & \text{if } \omega = 1 \end{cases} \quad (\text{A.5})$$

and D_L as the overall payoff of the strong faction following $d_{H,\omega_1} = 0$

$$D_L = \begin{cases} -g(1, 0) + \rho(\tilde{x}_1 = 1) (\beta R - \pi g(\phi_L, 1)) & \text{if } \omega = 0 \\ -g(1, 1) + \rho(\tilde{x}_1 = 1) (\beta R - \pi g(\phi_L, 1)) & \text{if } \omega = 1 \end{cases} \quad (\text{A.6})$$

First, notice that if parameters are such that $D_L(\rho(\tilde{x}_1 = 1) = 1) \leq E_L \forall \omega$, then, regardless of the value of $\rho(\tilde{x}_1 = 1)$, the faction chooses $d_1 = 1$.

Suppose instead that parameters are such that there exists at least an ω for which $D_L(\rho(\tilde{x}_1 = 1) = 1) > E_L$. Then, there exists a $\bar{\rho} \in [0, 1]$ such that, if off-path beliefs are such that $\rho(\tilde{x}_1 = 1) > \bar{\rho}$, then the strong faction prefers $d_1 = 0$. There is no such incentive for the weak faction. This, however, cannot be an equilibrium, because it would imply $\hat{\gamma}(1) < \gamma$, hence a contradiction with $\rho(1) > 0$. $\square$

$\square$

Proof of Lemma 3. First, note that, in any equilibrium with always dissent, the only relevant off-path event is $\tilde{x} = 1$. $\tilde{x} = 0$ may or may not be on path, but this does not affect the incentives of factions when choosing whether to dissent. Therefore, $\hat{\gamma}(1)$ cannot be pinned down by Bayes' rule and $\rho(1)$ can be anything.

However, in equilibrium $d_{H,\omega_1} = 1$ irrespective of $\rho(1)$. Consider now the strong faction. Consider parameters such that there exists at least an ω for which $D_L(\rho(\tilde{x}_1 = 1) = 1) > E_L$. In this case, then there exists a $\bar{\rho} \in [0, 1]$ such that, if off-path beliefs are such that $\rho(\tilde{x}_1 = 1) > \bar{\rho}$, then the strong faction has a profitable deviation to $d_1 = 0$. There is no such incentive for the weak faction. So, the only beliefs consistent with D1 imply that $\rho(\tilde{x}_1 = 1) = 0$, because any $\tilde{x}_1 = 1$

must come from the strong faction. Hence, as long as off-path beliefs satisfy D1, we have $d_1 = 1$ for all types and states. $\square$

Proof of Lemma 4. The proofs simply follows by observing that, as factions always dissent, $x_1 = 1$ fully reveals the type, and politicians care both about matching the state and remaining in power. Therefore, there is no trade off between these objectives in two cases. First, when $\sigma_{H,1} = 1$, the strong incumbent leadership is able to both match the state and gain re-election with certainty. Second, when $\sigma_{L,0} = 1$, the weak incumbent leadership is able to match the state, without revealing its type. $\square$

Proof of Lemma 5. With a slight abuse of notation, define $\hat{\gamma}(0)$ as the posterior probability of $\phi_I = \phi_H$ upon observing $\tilde{x}_1 = 0$. This is equal to:

$$\begin{aligned}\hat{\gamma}(0) \equiv \Pr(\phi_I = \phi_H | \tilde{x}_1 = 0) &= \frac{\Pr(\tilde{x}_1 = 0 | \phi_I = \phi_H) \gamma}{\Pr(\tilde{x}_1 = 0 | \phi_I = \phi_H) \gamma + \Pr(\tilde{x}_1 = 0 | \phi_I = \phi_L)(1 - \gamma)} \\ &= \frac{\sigma_{H,0} \gamma (1 - \pi)}{\sigma_{H,0} \gamma (1 - \pi) + (1 - \gamma)(1 - \pi + \pi(1 - \sigma_{L,1}))}.\end{aligned}$$

To see when $\sigma_{L,1} = 1$, notice that the voter re-elects the incumbent upon observing $\tilde{x}_1 = 0$, i.e., $\rho(0) > 0$, when:

$$\begin{aligned}\hat{\gamma}(0) &\geq \gamma \\ \sigma_{H,0}(1 - \pi) &\geq 1 - \pi + \pi(1 - \sigma_{L,1}) \\ 0 &\geq \pi(1 - \sigma_{L,1}) + (1 - \pi)(1 - \sigma_{H,0}).\end{aligned}\tag{A.7}$$

There could be two cases. Either there is an equilibrium where $\rho(0) > 0$, or $\rho(0) = 0$.

In the first case ($\rho(0) > 0$), it must be that $\sigma_{L,1} = \sigma_{H,0} = 1$, otherwise (A.7) never holds. In this equilibrium, the LHS and the RHS of (A.7) are equal and the voter is indifferent about whom to re-elect. Thus, in this case the weak incumbent leadership always matches the state.

In the second case ($\rho(0) = 0$), the weak incumbent leadership has no way to get re-elected, and always matches the state of the world to avoid the cost of mismatch. Therefore, regardless of the voter re-election rule, it is always optimal for the weak incumbent leadership to match the state. $\square$

Proof of Proposition 1. Let $\rho(0) = 1$. In such an equilibrium, the H type always matches the state and is re-elected, therefore there are no deviations from $\sigma_{H,\omega} = 1$ for any ω_1 . The same is true for the L type in state $\omega_1 = 0$. Therefore, we only need to check the incentive compatibility of the L type in state $\omega_1 = 1$. Its equilibrium payoff from choosing $x_1 = 1$ is

$$\mathbb{E}_{\omega_2}(u^p|\phi_L, \omega_1 = 1, x_1 = 1) = -f(\phi_L, 1),$$

because the implemented policy is ϕ_L and it is not re-elected. Deviating to $x_1 = 0$ yields:

$$\mathbb{E}_{\omega_2}(u^p|\phi_L, \omega_1 = 1, x_1 = 0) = -f(0, 1) + (1 - \beta)R - \pi f(\phi_L, 1).$$

Therefore, a full discipline equilibrium requires:

$$(1 - \beta)R \leq f(0, 1) - f(\phi_L, 1) + \pi f(\phi_L, 1).$$

Consider now $\rho(0) = 0$. This implies that the equilibrium strategy for $\phi_I = \phi_L$ is $\sigma_{L,\omega}^* = 1$ for every ω : since the incumbent leadership is never re-elected, it is always optimal to match the state.

Now let $\phi_I = \phi_H$. The incumbent compares the equilibrium payoff

$$\mathbb{E}_{\omega_2}(u^p|\phi_H, \omega_1 = 0, x_1 = 0) = 0$$

to the payoff from deviation

$$\mathbb{E}_{\omega_2}(u^p|\phi_H, \omega_1 = 0, x_1 = 1) = -f(\phi_H, 0) + (1 - \beta)R - \pi f(\phi_H, 1).$$

In equilibrium, $\sigma_{H,0}^* = 1$ iff

$$(1 - \beta)R \leq f(\phi_H, 0) + \pi f(\phi_H, 1).$$

To complete the proof, note that an equilibrium with $d_1 = 1 \forall \phi_I, \omega_1$ always exists, as long as off-path beliefs satisfy D1, as shown in Lemma 3.

To complete the proof, we need to show that an equilibrium with $d_1 = 1 \forall \phi_I, \omega_1$ always exists. We claim that, as long as off-path beliefs satisfy D1, such an equilibrium always exists. First, note that the only off-path event is $\tilde{x} = 1$. Therefore, $\hat{\gamma}(1)$ cannot be pinned down by Bayes' rule and $\rho(1)$ can be anything.

However, in equilibrium $d_{H,\omega_1} = 1$ irrespective of $\rho(1)$. Consider now the strong faction. Consider parameters such that there exists at least an ω for which $D_L(\rho(\tilde{x}_1 = 1) = 1) > E_L$. In this case, then there exists a $\bar{\rho} \in [0, 1]$ such that, if off-path beliefs are such that $\rho(\tilde{x}_1 = 1) > \bar{\rho}$, then the strong faction has a profitable deviation to $d_1 = 0$. There is no such incentive for the weak faction. So, the only beliefs consistent with D1 imply that $\rho(\tilde{x}_1 = 1) = 0$, because any $\tilde{x}_1 = 1$ must come from the strong faction. Hence, as long as off-path beliefs satisfy D1, we have $d_1 = 1$ for all types and states. $\square$

Proof of Proposition 2. Existence. Suppose $\rho(0) = 0$. This implies that the equilibrium strategy for $\phi_I = \phi_L$ is $\sigma_{L,1}^* = 1$: since the weak faction is never re-elected, it is optimal to at least match the state.

Now let $\phi_I = \phi_H$, and conjecture $\sigma_{H,0}^* = 0$. The incumbent compares the equilibrium payoff

$$\mathbb{E}_{\omega_2}(u^p | \phi_H, \omega_1 = 0, x_1 = 1) = -f(\phi_H, 0) + (1 - \beta)R - \pi f(\phi_H, 1),$$

to the payoff from deviation

$$\mathbb{E}_{\omega_2}(u^p | \phi_H, \omega_1 = 0, x_1 = 0) = 0.$$

In equilibrium, $\sigma_{H,0}^* = 0$ if and only if

$$(1 - \beta)R \geq f(\phi_H, 0) + \pi f(\phi_H, 1).$$

Finally, notice that given the equilibrium $\sigma_{I,\omega}$, it follows from Remark 1 that $\rho(0) = 0$.

To complete the proof, note that an equilibrium with $d_1 = 1 \forall \phi_I, \omega_1$ always exists, as long as off-path beliefs satisfy D1, as shown in Lemma 3.

Uniqueness. To show uniqueness, first notice that the equilibrium is either a full discipline or an over-reform one. Second, the full discipline equilibrium with $\rho(0) = 0$ exists iff $(1 - \beta)R \leq f(\phi_H, 0) + \pi f(\phi_H, 1)$. Third, the full discipline equilibrium with $\rho(0) = 1$ exists if and only if

$(1 - \beta)R \leq f(0, 1) - f(\phi_L, 1) + \pi f(\phi_L, 1)$. Therefore, the over-reform equilibrium is unique if and only if $(1 - \beta)R > \max\{f(\phi_H, 0) + \pi f(\phi_H, 1); f(0, 1) - f(\phi_L, 1) + \pi f(\phi_L, 1)\}$. $\square$

Proof of Proposition 3. We can express the voter welfare in the over-reform equilibrium as follows:

$$\begin{aligned} W &= \gamma [\pi (-f(\phi_H, 1) - \pi f(\phi_H, 1) - (1 - \pi)0) + (1 - \pi) (-f(\phi_H, 0) - \pi f(\phi_H, 1))] \\ &\quad + (1 - \gamma) [\pi (-f(\phi_L, 1) + W_2^C) + (1 - \pi)(0 + W_2^C)] \\ &= -2\gamma\pi f(\phi_H, 1) - \gamma(1 - \pi)f(\phi_H, 0) - (1 - \gamma)\pi f(\phi_L, 1) + (1 - \gamma)W_2^C \end{aligned} \quad (\text{A.8})$$

where W_2^C is the voter's expected second-period payoff if the challenger is in office, which is:

$$\begin{aligned} W_2^C &= \gamma [-\pi f(\phi_H, 1) - (1 - \pi)0] + (1 - \gamma) [-\pi f(\phi_L, 1) - (1 - \pi)0] \\ &= -\pi [\gamma f(\phi_H, 1) + (1 - \gamma)f(\phi_L, 1)]. \end{aligned}$$

Taking the derivative of W with respect to γ yields:

$$\frac{\partial W}{\partial \gamma} = -2\pi f(\phi_H, 1) - (1 - \pi)f(\phi_H, 0) + \pi f(\phi_L, 1) - W_2^C + \frac{\partial W_2^C}{\partial \gamma}(1 - \gamma),$$

which, substituting in the value of W_2^C , simplifies to:

$$\begin{aligned} \frac{\partial W}{\partial \gamma} &= -2\pi f(\phi_H, 1) - (1 - \pi)f(\phi_H, 0) + \pi f(\phi_L, 1) + \pi\gamma f(\phi_H, 1) \\ &\quad + \pi(1 - \gamma)f(\phi_L, 1) + \pi(1 - \gamma)f(\phi_L, 1) + \pi(1 - \gamma)f(\phi_H, 1) \\ &= (3 - 2\gamma)\pi [f(\phi_L, 1) - f(\phi_H, 1)] - (1 - \pi)f(\phi_H, 0) \end{aligned}$$

Thus, we can see that welfare is decreasing in γ if and only if:

$$\underbrace{(3 - 2\gamma)}_{\text{Benefit from selection}} \underbrace{\pi [f(\phi_L, 1) - f(\phi_H, 1)]}_{\text{Benefit from better reform}} < (1 - \pi)f(\phi_H, 0).$$

To prove the non-monotonicity in γ , note that the derivative is equal to zero when:

$$\begin{aligned}\pi(3 - 2\gamma) [f(\phi_L, 1) - f(\phi_H, 1)] - (1 - \pi)f(\phi_H, 0) &= 0 \\ \pi(3 - 2\gamma) [f(\phi_L, 1) - f(\phi_H, 1)] &= (1 - \pi)f(\phi_H, 0) \\ \gamma &= \frac{1}{2} \left[3 - \frac{(1 - \pi)f(\phi_H, 0)}{\pi(f(\phi_L, 1) - f(\phi_H, 1))} \right] := \gamma^*,\end{aligned}$$

where $(1 - \pi)f(\phi_H, 0)$ is the cost of a reform implemented in state 0 times the probability that that state is realized and $\pi(f(\phi_L, 1) - f(\phi_H, 1))$ is the benefit from having an H type, rather than a L type, implementing a reform in state 1.

To see that γ^* represents a maximum, we take the SOC:

$$\frac{\partial^2 W}{\partial \gamma^2} = -\pi [f(\phi_L, 1) - f(\phi_H, 1)] < 0.$$

It is immediate to see that

$$\begin{aligned}\gamma^* &> 0 \\ 3 &> \frac{(1 - \pi)f(\phi_H, 0)}{\pi(f(\phi_L, 1) - f(\phi_H, 1))}\end{aligned}$$

and

$$\begin{aligned}\gamma^* &< 1 \\ 1 &< \frac{(1 - \pi)f(\phi_H, 0)}{\pi(f(\phi_L, 1) - f(\phi_H, 1))}\end{aligned}$$

therefore, voter's welfare is non-monotonic in γ if and only if

$$\frac{(1 - \pi)f(\phi_H, 0)}{\pi(f(\phi_L, 1) - f(\phi_H, 1))} \in (1, 3), \tag{A.9}$$

which proves the statement. □

Proof of Proposition 4. For the full discipline equilibria, we need to consider two cases, based on the voter re-election rule when $\omega_1 = 0$.

First case: $\rho(0) = 0$. Let $W_{d,0}$ define the voter welfare in the full discipline equilibrium when $\rho(0) = 0$. We have:

$$\begin{aligned} W_{d,0} = & \gamma \left[\pi (-f(\phi_H, 1) - \pi f(\phi_H, 1)) + (1 - \pi) (0 + W_2^C) \right] \\ & + (1 - \gamma) \left[\pi (-f(\phi_L, 1) + W_2^C) + (1 - \pi)(0 + W_2^C) \right]. \end{aligned} \quad (\text{A.10})$$

Comparing (A.10) with (A.8), it is clear that welfare is higher in the over-reform equilibrium if and only if

$$\begin{aligned} \gamma(1 - \pi) (-f(\phi_H, 0) - \pi f(\phi_H, 1)) & > \gamma(1 - \pi)(0 + W_2^C) \\ -f(\phi_H, 0) - \pi f(\phi_H, 1) & > -\pi [\gamma f(\phi_H, 1) + (1 - \gamma)f(\phi_L, 1)]. \end{aligned}$$

Note that the RHS is maximized when $\gamma = 1$. Moreover, when $\gamma = 1$ the RHS is always higher than the LHS, therefore the condition never holds. Furthermore, the RHS is linearly increasing in γ . As $\gamma \rightarrow 0$, the RHS tends to $-\pi f(\phi_L, 1)$. Therefore, the condition holds for sufficiently low γ as long as $LHS > RHS(\gamma = 0)$. This simplifies to

$$\begin{aligned} -f(\phi_H, 0) - \pi f(\phi_H, 1) & > -\pi f(\phi_L, 1) \\ \pi[f(\phi_L, 1) - f(\phi_H, 1)] & > f(\phi_H, 0). \end{aligned}$$

Second case: $\rho(0) = 1$. We can express the voter welfare in the full discipline equilibrium as follows:

$$\begin{aligned} W_{d,1} = & \gamma [\pi (-f(\phi_H, 1) - \pi f(\phi_H, 1)) + (1 - \pi) (0 - \pi f(\phi_H, 1))] \\ & + (1 - \gamma) \left[\pi (-f(\phi_L, 1) + W_2^C) + (1 - \pi)(0 - \pi f(\phi_L, 1)) \right] \end{aligned} \quad (\text{A.11})$$

Comparing (A.11) with (A.8), it is clear that welfare is higher in the over-reform equilibrium iff

$$\begin{aligned} \gamma(1 - \pi)(-f(\phi_H, 0)) + (1 - \gamma)(1 - \pi)W_2^C & > \gamma(1 - \pi)0 - (1 - \gamma)(1 - \pi)\pi f(\phi_L, 1) \\ -\gamma f(\phi_H, 0) - (1 - \gamma)\pi [\gamma f(\phi_H, 1) + (1 - \gamma)f(\phi_L, 1)] & > -(1 - \gamma)\pi f(\phi_L, 1) \end{aligned}$$

$$\begin{aligned}
(1 - \gamma)\pi [\gamma f(\phi_L, 1) - \gamma f(\phi_H, 1)] &> \gamma f(\phi_H, 0) \\
(1 - \gamma)\pi [f(\phi_L, 1) - f(\phi_H, 1)] &> f(\phi_H, 0),
\end{aligned} \tag{A.12}$$

which never holds for sufficiently high γ . A sufficient condition for (A.12) to be satisfied at a low γ is that

$$\pi [f(\phi_L, 1) - f(\phi_H, 1)] > f(\phi_H, 0),$$

which completes the proof. $\square$

Proof of Lemma 6. Conjecture an equilibrium such that $d_1 = 1$ for $\phi_I = \phi_H$ and $d_1 = 0$ for $\phi_I = \phi_L$. This implies $\rho(1) = 0$, since it is only the strong faction that dissents. Consider the incentives of the strong faction ($\phi_I = \phi_L$), when $\omega_1 = 1$ and $x_1 = 1$. The equilibrium payoff (from $d_1 = 0$) is:

$$-g(1, 1) + \rho(1) [\beta R - \pi g(\phi_L, 1)], \tag{A.13}$$

whereas a deviation to $d_1 = 1$ yields:

$$-g(\phi_L, 1) + \rho(\phi_L) [\delta \beta R - \pi g(\phi_L, 1)] \tag{A.14}$$

Because of the voter retention rule in equilibrium, (A.14) is always higher than (A.13), which implies a contradiction. $\square$

Proof of Proposition 5. Suppose $\rho(1) = 1$. The following are the incentive compatibility conditions for $d_{1,H} = 0$ and $d_{1,L} = 0$, respectively:

$$(1 - \delta)\beta R \geq g(1, \omega) - g(\phi_H, \omega) \tag{A.15}$$

$$\beta R \geq g(1, \omega) - (1 - \pi)g(\phi_L, \omega), \tag{A.16}$$

which can be expressed as

$$\beta R \geq \max \left[\frac{g(1, \omega) - g(\phi_H, \omega)}{1 - \delta}; g(1, \omega) - (1 - \pi)g(\phi_L, \omega) \right], \tag{A.17}$$

for $\omega \in \{0, 1\}$.

Given that factions do not dissent, leaders face the following incentives: For the weak leader ($\phi_I = \phi_L$):

- When $\omega_1 = 1$: Choosing $x_1 = 1$ gives payoff $-f(1, 1) + (1 - \beta)R - \pi f(\phi_L, 1)$. Choosing $x_1 = 0$ gives payoff $-f(0, 1)$.
- When $\omega_1 = 0$: Choosing $x_1 = 1$ gives payoff $-f(1, 0) + (1 - \beta)R - \pi f(\phi_L, 1)$. Choosing $x_1 = 0$ gives payoff 0.

Thus, the weak leader matches the state if and only if:

$$(1 - \beta)R \leq f(1, 0) + \pi f(\phi_L, 1). \quad (\text{A.18})$$

For the strong leader ($\phi_I = \phi_H$): when $\omega_1 = 0$, he matches the state. When $\omega_1 = 1$, choosing $x_1 = 1$ (over-reform) gives payoff $-f(1, 1) + (1 - \beta)R - \pi f(\phi_H, 1)$. Choosing $x_1 = 0$ gives a payoff of 0. The strong leader will then over-reform if and only if:

$$(1 - \beta)R \geq f(1, 0) + \pi f(\phi_H, 1). \quad (\text{A.19})$$

Given the proposed strategies, the voter's beliefs are $\Pr(\phi_I = \phi_H | \tilde{x}_1 = 1) = 1$ (only strong leaders implement reforms), $\Pr(\phi_I = \phi_H | \tilde{x}_1 = 0) = 0$ (only weak leaders choose no reform). Therefore, $\rho(1) = 1$ and $\rho(0) = 0$ are sequentially rational. $\square$

Proof of Proposition 6. Let us first define the voter's ex ante welfare in the over-reform equilibrium with active factions (W) and in the over-reform equilibrium with inactive factions in period 1 (W_{noF}), recalling that they are active in period 2.

$$\begin{aligned} W &= \gamma \left[\pi \left(-f(\phi_H, 1) + V_H \right) + (1 - \pi) \left(-f(\phi_H, 0) + V_H \right) \right] \\ &\quad + (1 - \gamma) \left[\pi \left(-f(\phi_L, 1) + V_C \right) + (1 - \pi) \left(0 + V_C \right) \right], \\ W_{noF} &= \gamma \left[\pi \left(0 + V_H \right) + (1 - \pi) \left(-f(1, 0) + V_H \right) \right] \\ &\quad + (1 - \gamma) \left[\pi \left(0 + V_L \right) + (1 - \pi) \left(0 + V_C \right) \right]. \end{aligned}$$

We can now calculate the difference as:

$$\begin{aligned}
W - W_{noF} &= \gamma \pi \left(-f(\phi_H, 1) - 0 \right) + \gamma (1 - \pi) \left(-f(\phi_H, 0) + f(1, 0) \right) \\
&\quad + (1 - \gamma) \pi \left(-f(\phi_L, 1) - 0 \right) + (1 - \gamma) \pi \left(V_C - V_L \right).
\end{aligned}$$

where V_H , V_L and V_C is the expected second period welfare if the politician in power in period 2 is of high type, low type or is a challenger.

Since

$$V_C - V_L = [\gamma V_H + (1 - \gamma) V_L] - V_L = \gamma (V_H - V_L) = \gamma \pi (-f(\phi_H, 1) + f(\phi_L, 1))$$

one can substitute to get

$$\begin{aligned}
W - W_{noF} &= \gamma \pi \left(-f(\phi_H, 1) \right) + \gamma (1 - \pi) \left(-f(\phi_H, 0) + f(1, 0) \right) \\
&\quad + (1 - \gamma) \pi \left(-f(\phi_L, 1) \right) + (1 - \gamma) \pi \gamma \pi \left(-f(\phi_H, 1) + f(\phi_L, 1) \right).
\end{aligned}$$

We can further simplify and collect terms to capture the various components of the comparison:

$$\begin{aligned}
W - W_{noF} &= \underbrace{\gamma \pi \left(-f(\phi_H, 1) \right) + (1 - \gamma) \pi \left(-f(\phi_L, 1) \right)}_{\text{cost of policy attenuation}} \\
&\quad + \underbrace{(1 - \gamma) \gamma \pi^2 \left(f(\phi_L, 1) - f(\phi_H, 1) \right)}_{\text{benefit on selection}} \\
&\quad + \underbrace{\gamma (1 - \pi) \left(f(1, 0) - f(\phi_H, 0) \right)}_{\text{benefit of less damaging over-reforms}}.
\end{aligned}$$

It is straightforward to see that, in the limit for $\pi \rightarrow 0$ or for $\gamma \rightarrow 1$, the condition is met. $\square$

B. Extension: Probabilistic type revelation

Suppose that $x_1 = 1$ reveals the type of the incumbent only with probability q . In other words, the voter observes $\tilde{x}_1$ with probability q , and does not learn anything with probability $1 - q$. We look for conditions that guarantee the existence of an over-reform equilibrium, where $\sigma_{H,0} = 0$ and $\sigma = 1$ otherwise.

First, suppose the voter observes the policy. In this case, the re-election probability is the same as in the main body of the paper. Second, suppose the voter doesn't learn anything. In this case, he is indifferent between re-electing or not, and therefore any re-election probability $\rho(\emptyset)$ is sequentially rational.

Given this, the high type chooses $x_1 = 1$ when $\omega_1 = 0$ iff

$$-f(\phi_H, 0) + q(R - \pi f(\phi_H, 1)) + (1 - q)\rho(\emptyset)(R - \pi f(\phi_H, 1)) \geq (1 - q)\rho(\emptyset)(R - \pi f(\phi_H, 1)),$$

which simplifies to

$$q > \frac{f(\phi_H, 0)}{R - \pi f(\phi_H, 1)}. \quad (\text{B.1})$$

The incentive structure for the low type is the same as before: if $x_1 = \omega_1$, he's re-elected with probability $(1 - q)\rho(\emptyset)$, and pays no cost of policy mismatch. If $x_1 \neq \omega_1$, he pays the cost of policy mismatch without being re-elected. Hence, in equilibrium the low type always chooses to match the state of the world.

Therefore, the Over-Reform equilibrium exists as long as Condition (B.1) is satisfied.

C. Extension: different degrees of factionalization

In this appendix, we consider a variation of the benchmark model where we allow the prior on factionalization to be different between challenger and incumbent. More formally, we assume

$$Pr(\phi_I = \phi_H) = \gamma^I \neq \gamma^C = Pr(\phi_C = \phi_H).$$

The rest of the game is unchanged.

C.1. Equilibria

First, we establish that the over-reform equilibrium exists under the same conditions, irrespective of γ^I and γ^C .

Proposition C1. *There exists an “over-reform” equilibrium where*

- (i) $\sigma_{H,0} = 0$ and $\sigma_{I,\omega} = 1$ otherwise;
- (ii) $\rho(0) = 0$; if and only if $R \geq f(\phi_H, 0) + \pi f(\phi_H, 1)$.

Proof of Proposition C1. In the over-reform equilibrium, $Pr(\phi_I = \phi_H | \tilde{x}_1 = 0) = 0$, because $\sigma_{H,0} = 1$. This was the only condition in Proposition 2 affected by γ , and it does not change when $\gamma^I \neq \gamma^C$. Applying the same logic as in the proof of Proposition 2 completes this proof. $\square$

Therefore, the over-reform equilibrium is fully robust to this variation in the modelling assumptions. Other pure strategy equilibria may or may not exist, depending on the ranking between γ^I and γ^C .

First, consider $\gamma^I < \gamma^C$. In this case, only the full discipline equilibrium with $\rho(0) = 0$ survives.

Proposition C2. *If $\gamma^I < \gamma^C$, there exists a full discipline equilibrium where*

- (i) $\sigma_{I,\omega} = 1$ for every ω, ϕ_I , and
- (ii) $\rho(0) = 0$ if and only if $R \leq f(\phi_H, 0) + \pi f(\phi_H, 1)$. There are no other pure strategy full discipline equilibria.

Proof of Proposition C2. First, note that Lemma 4 applies to this case as well. Second, note that

$$\hat{\gamma}^I(0) = \frac{\sigma_{H,0}\gamma^I(1-\pi)}{\sigma_{H,0}\gamma^I(1-\pi) + (1-\gamma^I)(1-\pi + \pi(1-\sigma_{L,1}))}$$

is increasing in both $\sigma_{H,0}$ and $\sigma_{L,1}$. By substitution, $\max[\hat{\gamma}^I(0)] = \gamma^I < \gamma^C$, therefore in every equilibrium it must be that $\rho(0) = 0$. The rest follows from the proof of Proposition 1. $\square$

Note that there are no other equilibria, in this case.

Second, consider $\gamma^I > \gamma^C$. In this case, only the full discipline equilibrium with $\rho(0) = 1$ survives. On top of this, under some conditions it is possible to have an under-reform equilibrium, where the strong leader chooses the correct policy and the weak leader never implements a reform.

Proposition C3. *If $\gamma^I > \gamma^C$, there exists a full discipline equilibrium where*

- (i) $\sigma_{I,\omega} = 1$ for every ω, ϕ_I , and
- (ii) $\rho(0) = 1$ if and only if $R \leq f(0, 1) - f(\phi_L, 1) + \pi f(\phi_L, 1)$;
There are no other pure strategy full discipline equilibria.

Proof of Proposition C3. First, note that Lemma 4 applies to this case as well. Second, note that in a full discipline equilibrium

$$\hat{\gamma}^I(0) = \gamma^I > \gamma^C$$

therefore in every equilibrium it must be that $\rho(0) = 1$. The rest follows from the proof of Proposition 1. $\square$

Proposition C4. *If $\gamma^I > \gamma^C$, there exists an “under-reform” equilibrium where*

(i) $\sigma_{L,1} = 0$ and $\sigma_{I,\omega} = 1$ otherwise, and

(ii) $\rho(0) = 1$ if and only if $R \geq f(0,1) - f(\phi_L,1) + \pi f(\phi_L,1)$ and $\gamma^I \geq \frac{\gamma^C}{1-\pi(1-\gamma^C)}$.

Proof of Proposition C4. First, note that Lemma 4 applies to this case as well. Second, note that if $\rho(0) = 1$, then in equilibrium it must be that $\sigma_{H,0} = 1$ because there are policy gains without losses in terms of re-election chances. Third, note from the proof of Proposition 1 that if $\rho(0) = 1$ and $R \geq f(0,1) - f(\phi_L,1) + \pi f(\phi_L,1)$, then in equilibrium it must be that $\sigma_{L,1} = 0$. Otherwise, $\sigma_{L,1} = 1$ and we fall into the full discipline equilibrium described above. Finally, for this equilibrium to exist, it must be that $\rho(0) = 1$. With the strategies outlined above, this implies

$$\begin{aligned}\hat{\gamma}^I(0) &= \frac{\gamma^I(1-\pi)}{\gamma^I(1-\pi) + 1 - \gamma^I} \geq \gamma^C \\ \frac{\gamma^I(1-\pi)}{1-\pi\gamma^I} &\geq \gamma^C \\ \gamma^I &\geq \frac{\gamma^C}{1-\pi(1-\gamma^C)}\end{aligned}$$

Note that if this condition is violated, $\rho(0)$ cannot be 1 and therefore the equilibrium does not exist. $\square$

As $1 - \pi(1 - \gamma^C) < 1$, Proposition C4 implies that γ^I must be sufficiently larger than γ^C for this equilibrium to exist. Secondly, as this equilibrium requires a sufficiently high R , it may co-exist with the over-reform one.

C.2. Factionalization and welfare

Having established that the over-reform equilibrium keeps existing under the same conditions, we can now study the effect of γ^I and γ^C on voter’s welfare separately. Therefore, rather than asking how aggregate factionalization affects welfare, we separate between the factionalization of the incumbent and of the challenger.

Proposition C5. *In the over-reform equilibrium, voter welfare (W) can be increasing or decreasing in factionalization of the incumbent (γ^I). Voter welfare is decreasing in factionalization of the challenger.*

Proof of Proposition C5. We can express the voter welfare in the over-reform equilibrium as follows:

$$\begin{aligned}
W &= \gamma^I [\pi (-f(\phi_H, 1) - \pi f(\phi_H, 1) - (1 - \pi)0) + (1 - \pi) (-f(\phi_H, 0) - \pi f(\phi_H, 1))] \\
&\quad + (1 - \gamma^I) [\pi (-f(\phi_L, 1) + W_2^C) + (1 - \pi)(0 + W_2^C)] \\
&= -2\gamma^I \pi f(\phi_H, 1) - \gamma^I (1 - \pi) f(\phi_H, 0) - (1 - \gamma^I) \pi f(\phi_L, 1) + (1 - \gamma^I) W_2^C
\end{aligned}$$

where W_2^C is the voter's expected second-period payoff if the challenger is in office, which is:

$$\begin{aligned}
W_2^C &= \gamma^C [-\pi f(\phi_H, 1) - (1 - \pi)0] + (1 - \gamma^C) [-\pi f(\phi_L, 1) - (1 - \pi)0] \\
&= -\pi [\gamma^C f(\phi_H, 1) + (1 - \gamma^C) f(\phi_L, 1)].
\end{aligned}$$

Taking the derivative of W with respect to γ^I yields:

$$\begin{aligned}
\frac{\partial W}{\partial \gamma^I} &= -2\pi f(\phi_H, 1) + \pi f(\phi_L, 1) - W_2^C - (1 - \pi) f(\phi_H, 0) \\
&= \pi [f(\phi_L, 1) - f(\phi_H, 1)] + \pi [\gamma^C f(\phi_H, 1) + (1 - \gamma^C) f(\phi_L, 1)] - (1 - \pi) f(\phi_H, 0) - \pi f(\phi_H, 1) \\
&= \pi (2 - \gamma^C) [f(\phi_L, 1) - f(\phi_H, 1)] - (1 - \pi) f(\phi_H, 0)
\end{aligned}$$

Thus, we can see that welfare is decreasing in γ if and only

$$\pi (2 - \gamma^C) [f(\phi_L, 1) - f(\phi_H, 1)] < (1 - \pi) f(\phi_H, 0).$$

Finally, note that

$$\text{sign} \left(\frac{\partial W}{\partial \gamma^C} \right) = \text{sign} \left(\frac{\partial W_2^C}{\partial \gamma^C} \right) = \pi [f(\phi_L, 1) - f(\phi_H, 1)] > 0,$$

therefore W is increasing in γ^C . □

Proposition C5 and its proof reveal several interesting results. First, factionalization of the challenger is always bad for voter's welfare. This is, however, a feature of the two-period structure of the model. As in period 2 everyone chooses the correct action, the stronger is the leadership of the challenger the better it is in case of a challenger's victory.

Second, the trade-off between over-reform and better implementation remains in this case as well. In fact, W is decreasing in γ^I (therefore, a more factionalized incumbent can be welfare improving) when

$$\pi(2 - \gamma^C) [f(\phi_L, 1) - f(\phi_H, 1)] < (1 - \pi)f(\phi_H, 0). \quad (\text{C.1})$$

Condition (C.1) is more likely to be satisfied if the expected cost of the over-reform (i.e., $(1 - \pi)f(\phi_H, 0)$) is high, if the expected benefit from correctly implementing a needed reform (i.e., $\pi[f(\phi_L, 1) - f(\phi_H, 1)]$) is low and if γ^C is high, i.e. if it is easy to replace a strong incumbent with a strong challenger.

D. Extension: Continuous Dissent Choice

In the main body of the paper we assume for simplicity that the faction either fully dissents, therefore inducing $\tilde{x}_t = \phi_I x_t$, or does not, therefore inducing $\tilde{x}_t = x_t$. In practice, factions can modulate their approval of the policies chosen by the party.

To capture this, in this Appendix we allow the dissent choice to be continuous. Relaxing this assumption is particularly relevant to show that the faction's equilibrium behavior of 'always dissent' does not depend on the fact that the strong faction cannot mimic the weak one, which is a feature of our baseline where dissent is a dichotomous choice.

To preserve the same notation of the baseline model, we assume that the faction decides how much to accommodate (a_t) the policy chosen by the leader, where $a_t \in [\phi_I, 1]$, with $0 < \phi_L < \phi_H < 1$. The implemented policy is then $\tilde{x}_t = a_t x_t$.

The timing of the game is as follows:

1. First period

- The incumbent observes ω_1 and chooses x_1 .
- The faction observes ω_1 and x_1 and chooses a_1 .
- The voter observes $\tilde{x}_1$, updates on γ , and votes.

2. Second period

- The incumbent observes ω_2 and chooses x_2 .
- The faction observes ω_2 and x_2 , and chooses a_2 .
- The policy is implemented.

We are interested in conditions such that both types of factions ‘fully’ dissent (i.e., up to the maximum possible to both) in equilibrium.

Proposition D1. *There exists an over-reform equilibrium where $a_t = \phi_I$ for $t = 1, 2$.*

Proof of Proposition D1. Conjecture an over-reform equilibrium, where $a_t = \phi_I$. We assume that, upon observing the off-path event $\tilde{x}_t \in (\phi_L, \phi_H)$, the voter ousts the incumbent, while the voter re-elects with probability η if $\tilde{x}_t \in (\phi_H, 1]$.

For the weak faction, the analysis is straightforward: there is never a profitable deviation from $a_t = \phi_H$, since it signals that the leader is strong (therefore securing re-election) with the minimum policy loss.

Hence, we focus on the decision of the strong faction ($\phi_I = \phi_L$). In the second period, the incentives are straightforward: in the absence of re-election concerns, the strong faction always chooses $a_2 = \phi_L$ because it minimizes the policy loss.

In the first period, the best deviation for the strong faction is $a_1 = \phi_H$, guaranteeing re-election at the lowest policy cost. We now check for deviations to this strategy.

Suppose $x_1 = 1$ and $\omega_1 = 1$. There exists no profitable deviation from $a_1 = \phi_L$ if and only if:

$$-g(\phi_L, 1) \geq -g(\phi_H, 1) + \beta R - \pi g(\phi_L, 1). \quad (\text{D.1})$$

Now let $\omega_1 = 0$. There exists no profitable deviation from $a_1 = \phi_L$ if and only if:

$$-g(\phi_L, 0) \geq -g(\phi_H, 0) + \beta R - \pi g(\phi_L, 1). \quad (\text{D.2})$$

Note that, as long as $a_t = \phi_I$, the conditions sustaining an over-reform equilibrium remain the same. Therefore, there exist values of βR sufficiently small such that the over-reform equilibrium exists. $\square$

The conditions (D.1 and D.2) are intuitive: as long as the faction's share of rents is sufficiently small, dissent is not too costly even if it implies losing the next election. Suppose instead the faction obtains a considerable share of rents: in this case, the over-reform equilibrium becomes harder to sustain for two reasons. First, the leadership is less willing to pay the cost of over-reforming today, obtaining a smaller share of rents tomorrow. Second, the faction is more motivated to pursue re-election, and therefore more likely to choose a milder opposition.