

Embracing the Enemy*

Álvaro Delgado-Vega[†] Johannes Schneider[‡]

July 2026

Abstract

A principal (such as a centrist political party) can partially influence the allocation of power between two competing parties. The principal is closer to one party, the “Friend”, than to the other, the “Enemy”. The principal’s optimal contract initially seeks to exclude the Enemy. However, once the Enemy gains power, the principal embraces him in exchange for policy moderation. Moderation also disciplines the Friend, inducing him to move closer to the principal’s preferred policy. Principals close to the Friend fully embrace the Enemy; more centrist principals divide their support. Commitment benefits the principal only if she is close to the Friend and parties value power little.

JEL CODES: D23, D74, D86

KEYWORDS: Relational Contracts, Organizational Economics, Political Economy, Horizontal Moral Hazard, Coalition Formation

Politics is the art of the possible, the attainable—the art of the next best.

Otto von Bismarck

1 Introduction

All organizations must grapple with disagreement, especially when a discordant actor is in a position to shape collective decisions. One common strategy for limiting such influence is exclusion. Economists formalize this logic as the *Ally Principle*: whenever possible, the principal delegates authority to aligned actors, thereby preventing the discordant actor

*We are grateful to Antoine Loeper for inspiring discussions early on. We thank Antonio Cabrales, Paweł Doligalski, Dana Foarta, Gabriele Gratton, Elliot Lipnowski, Chara Papioti, Spencer Pantoja, Harry Pei, Elia Sartori, Christoph Schottmüller, Jonathan B. Slapin, Takuo Sugaya, Heidi Thysen, Jan Zapal, and audiences at various places. The editor and three anonymous referees provided excellent comments that shaped the final form of the paper. We are indebted to them. Johannes Schneider acknowledges financial support from PID-2024-149885OB-I00. We thank the institutions we spent time at during this research for their hospitality and inspiration: Schneider was at uc3m, UBern, and NHH (Bergen); Delgado-Vega was at ETH Zurich and U. Carlos III de Madrid.

[†]University of Chicago; adelgadovega@uchicago.edu

[‡]Johannes Kepler University; johannes.schneider@jku.at

from exercising power. Its political counterpart, the *cordon sanitaire*, follows the same logic: mainstream parties collaborate to isolate radical parties and deny them meaningful influence.

Yet permanent exclusion proves challenging in practice. Consider, for instance, a parliament witnessing the emergence of an anti-establishment party. This newcomer competes directly with established parties, disrupting traditional governance patterns and reshaping the distribution of political power. Centrist parties—strategically positioned between ideological extremes—may become power brokers deciding, for example, who obtains cabinet posts. Their choices, however, are constrained: parliamentary arithmetic or shifts in public sentiment limit their options. A similar problem arises for centrist interest groups allocating electoral support—think of campaign contributors, major newspapers, or even groups of (perhaps organized) voters.¹ By endorsing aligned parties, they may bring them to power more often, but they cannot avoid the risk that a different party ultimately wins office. In this paper, we study optimal power brokerage in such scenarios. Our power broker, though influential, is not all-powerful. What strategy should she pursue? Should she impose a *cordon sanitaire* around a discordant agent? Or should she, on the contrary, back him in exchange for his moderation?

We consider an infinite-horizon setting in which two political parties (both “he”) repeatedly compete for the right to set policy. Both are power-hungry and agenda-driven, vying over *who* leads and *which* policy is made. A third player, the principal (“she”), is the power broker. She has an agenda, too, and only cares about that. The principal cannot choose the policy, but influences who gets the right to choose. She can increase a party’s chances of obtaining that right, yet only to a limited extent. The principal’s ideal policy is between the ideals of the parties but closer to one (the *Friend*) than the other (the *Enemy*).

We study the dynamics of the principal’s optimal contract and derive three predictions: (i) at first, setting up a *cordon sanitaire* around the *Enemy* is optimal; (ii) eventually, that *cordon sanitaire* will fall; and (iii) once it falls, it is not optimal to restore it but to embrace the *Enemy*.

Our analysis unveils an inherently dynamic logic. If the game were played only once, whichever party received the right to act (the *leading party*) would simply choose his preferred policy. Anticipating this polarization, the principal would *fully endorse* the *Friend*, that is, she would do everything in her power to maximize the *Friend*’s likelihood of taking the lead. However, if the interaction occurs repeatedly, fully endorsing the *Friend* severely undermines the principal’s influence. To see why, suppose the *Enemy* becomes the leading party today. The principal can offer him the following contract: “If, today,

¹Voters and political parties can be thought of as having relational contracts that, even if sketchy, may be fully understood and relevant for politicians’ decision-making. Although some features of our equilibrium are complex and may require too much from a voter, we cannot rule it out; we leave it to the reader to interpret in each application whether traces of our results can be observed in voters’ behavior.

you choose a policy close to my bliss point, I will endorse you tomorrow and ever after.”²

We show that this exchange—trading future endorsement for present policy moderation—is effective in disciplining the Enemy but especially relevant if the Friend and the principal are closely aligned. Suppose, for example, that the Friend and the principal have identical preferences over policy and assume the principal has commitment power. In that case, the principal endorses the Friend, upholding the cordon sanitaire, as long as the Enemy is kept from leading. But once the Enemy leads for the first time, the principal flips and persistently embraces the Enemy in exchange for his moderation. Meanwhile, the Friend continues to implement the principal’s preferred policy whenever he leads.

If the principal and the Friend are not fully aligned, the principal faces another challenge: to incentivize the Friend to moderate his policy, too. Naturally, the principal can promise to endorse the Friend in exchange. But then, her embrace of the Enemy decreases, and the Enemy moderates less. For that reason, the principal prefers to incentivize the Friend indirectly: she leverages the fact that the Friend benefits from the Enemy’s moderation, too. By moderating the Enemy whenever he is in charge, the principal improves the Friend’s expected payoffs. Threatening to remove that moderation is enough to get the Friend to choose the principal’s preferred policy. Yet, if they disagree more, direct endorsement of the Friend becomes necessary.

Principals with more centrist bliss points have higher ex-ante payoffs if they can incentivize the Friend—either directly or indirectly—to choose their preferred policy. Intuitively, if the principal can drag the Friend’s policy choice to her bliss point, she is mechanically better off the closer she is to the Enemy. However, if the principal is too distant from the Friend, she cannot count on a fully complicit Friend. Then, being more centrist makes her worse off: neither the Friend’s nor the Enemy’s policy can be pulled more to the center. Yet, the principal is further from the Friend, whose policy choices are more relevant because of the initial cordon sanitaire.

A natural concern is the principal’s ability to uphold her promises. To address this, we devote the second part of the paper to the case in which the principal lacks *commitment power*, so all promises need to be self-enforcing. Commitment power has two advantages for the principal. First, in the optimal contract, the principal promises to endorse the Enemy in future periods in exchange for his moderation today. With commitment, this promise is trivially *credible*. Second, a committed principal can *credibly threaten* to punish a deviating Friend. Both threat and promise are curtailed if the principal lacks commitment because she genuinely aligns more with the Friend and therefore, irrespective of her endorsement, always hopes for the Friend to take the lead next.

We show that in some parameter regions, whether the principal can indeed commit bears no practical relevance. We can replicate the commitment outcome if (i) agents

²This contract, of course, hinges on the principal’s ability to credibly promise that endorsement. Whether and when she can do that is a core aspect we will discuss in what follows.

attach sufficient intrinsic value to leading or (ii) the principal is sufficiently balanced. Does cooperation unravel completely otherwise? We show that this depends on the level of alignment between principal and Friend. If the principal and the Friend fully align, the results are drastic: either the principal can implement the commitment outcome, or she cannot make use of dynamic incentives at all. Outside this extreme case, the unraveling is gradual. Even if the commitment contract cannot be implemented, the principal can often promise some endorsement to the Enemy, thereby maintaining some moderation.

Our predictions relate to the Ally Principle mentioned in the beginning (see also Bendor et al., 2001; Bendor and Meirowitz, 2004; Callander, 2008). In our model, the Ally Principle is not invalidated in general, but should be abandoned as soon as the Enemy gains agency. Thus, we explain *when* the Ally Principle holds, and when and why it gets discarded. Indeed, the behavior we outlined resembles what we can observe when parliamentary democracies deal with a new extremist party. At first, mainstream parties try to exclude it through a cordon sanitaire. However, once the newcomer attains agency, they switch to an embracing strategy to moderate it. Often, the extremist then integrates into the political establishment. As examples, we document this pattern both in Weimar Germany and postwar Italy.

One may wonder if the principal benefits from the existence of a competitive setting. Absent the Enemy, the principal would lose the disciplining effect of the power struggle on the Friend, but also dispose of the Enemy’s undesired policy choices. The Friend would choose the policy always, but the principal would lack any means of pulling his choice toward the center. We show that if the principal is non-extreme in her relative position and interactions are frequent enough, the benefits of competition are larger.

The paper proceeds as follows. After a brief review of the literature, we present the model in Section 2. Section 3 analyzes a special case, and Section 4 extends the analysis to the general case. Section 5 discusses the implications of the results and their connection to real-world cases. We conclude in Section 6. Proofs are in the Appendix.

Related Literature. Allocation of power is central in organizational economics and political economy (see Bolton and Dewatripont, 2013, for a survey). We focus on repeated interactions. Much of that literature considers settings where players fully control power allocation (e.g., Li et al., 2017; Lipnowski and Ramos, 2020; Acharya et al., 2024; Luo et al., 2025) or where that allocation is stochastic and exogenous (e.g., Dixit et al., 2000). We adopt a middle ground; our principal partly influences the power allocation.

To our knowledge, Delgado-Vega (2025) is the only paper with a similar setup. However, while that model hinges on a common resource pool—which agents use either to provide public goods or to bribe the principal for private gain—our model focuses on agents in pure conflict over a fixed pie. Moreover, the principal’s role differs: In Delgado-Vega (2025) the principal extracts common resources for personal benefit. Our principal is solely concerned with distribution. She leverages the agents’ conflict to implement her preferred

allocation. Conceptually, that means we face—unlike in the literature cited above—a setting of horizontal moral hazard: shirking is not the main incentive problem, but the direction of the action is. Consequently, the differentiation between Friend and Enemy becomes a key feature of our model.

Our model connects to the relational contracting literature (see, e.g., the surveys of Malcomson, 2012; MacLeod, 2014; Watson, 2021). Other papers within that literature consider relationships with multiple competing agents.³ The closest among these are Board (2011), Barron and Powell (2019), and Calzolari and Spagnolo (2020), who, like us, study direct and persistent competition between agents. In contrast to these papers, we analyze a non-transferable utility (NTU) setting that precludes front-loading payments. This NTU framework and our view of organizations as communities of fate—with persistent internal disagreements—connect our paper to the delegation literature (e.g., Dessein, 2002; Alonso et al., 2008; Alonso and Matouschek, 2008). Different from us, that delegation literature focuses on information use. In our model of complete and perfect information, delegation has other inherent frictions: the principal is not only unable to censor the agents’ action spaces, but also only partially controls *which* agent to entrust with power. We explore how the principal optimally uses that limited delegation power.

The idea of trading power for policy is central to both political economy (see, e.g., Anesi and Buisseret, 2022; Acharya et al., 2024; Invernizzi and Ting, 2024) and organizational economics (see, e.g., Gibbons and Murphy, 1992; Li et al., 2017; Gibbons, 2020). We contribute by studying relational incentives in a setting of persistent partisan disagreement and stochastically shifting authority.

Our setting is a repeated game with the particular feature that a subset of the players (the two parties) engages in a constant-sum game while a third player (the principal) acts as a power broker to her own advantage. As we will see, non-stationary optimal contracts appear because this power broker lacks full control over delegation. Moreover, the game is asymmetric because the principal is more closely aligned with one party, the Friend. We show that, in this setting, the principal’s value for commitment depends non-trivially on her relative position. This affects both her ability to punish the Friend after a deviation and the credibility of her promises to embrace the Enemy on the equilibrium path. Unlike Abreu (1986), our punishments combine grim trigger (always for the Enemy, sometimes for the principal) and sticks-and-carrots (for the Friend). The carrot is returning to the on-path continuation contract. The constant-sum property allows us to derive penal codes and thereby characterize the equilibrium payoff set.⁴

³Examples include Levin (2002), Rayo (2007), Board (2011), Andrews and Barron (2016), Barron and Powell (2019), Calzolari and Spagnolo (2020), and Nocke and Strausz (2023).

⁴In a stochastic game, Deb et al. (2025) study a principal who, like ours, exploits the constant-sum competition between two agents. However, their research question, model, and analysis differ substantially.

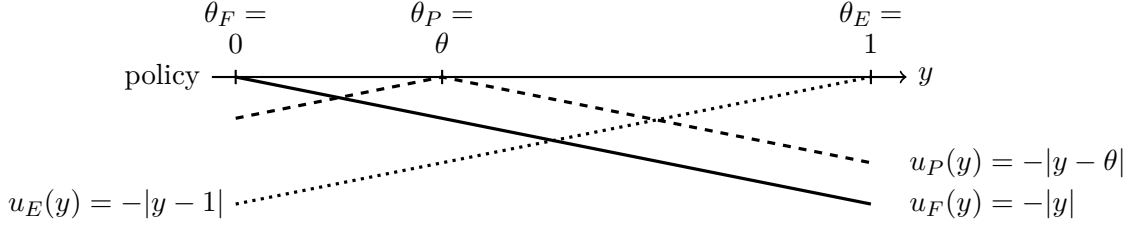


Figure 1: Bliss points and policy payoffs.

2 Model

Time is discrete and indexed by $t = 1, 2, \dots$. There are three risk-neutral players: one principal, P ; and two parties, the Friend, F , and the Enemy, E . Players discount the future exponentially with discount factor $\beta \in (0, 1)$. Throughout, information is perfect, and our players engage in the following repeated game.⁵

Stage Game. At the beginning of every period t , P provides a level of endorsement $s_t \in [-m, m]$, where $s_t = -m$ implies full endorsement of F , and $s_t = m$ implies full endorsement of E . Then, nature tosses a biased coin that yields $k_t \in \{F, E\}$. Endorsement s_t affects the probabilities of this biased coin. The probability of $k_t = E$ is $p(s_t) = 1/2 + s_t$, and $k_t = F$ realizes with the complementary probability, $1 - p(s_t)$. The principal's power is limited—that is, $m \in (0, 1/2)$. If k realizes, we say party k is the *leading party*. The leading party chooses a policy $y_t \in [0, 1]$. The stage game ends and all players collect their within-period payoffs.

Flow Payoffs. If the leader's policy is y , player $i \in \{F, E, P\}$ receives a stage payoff from that policy given by

$$u_{i,t}(y) := -|y - \theta_i|,$$

where θ_i denotes player i 's persistent bliss point in the policy space,

$$\theta_F \equiv 0 \quad \theta_P \equiv \theta \in [0, 1/2) \quad \theta_E \equiv 1.$$

Only the party who leads in the current period, k , receives the leadership rent, $b > 0$, in addition to $u_{k,t}$. Notationally, it is useful to distinguish between a player's continuation payoff at the beginning of a period, $w_i(\cdot)$, and *after selection*, $v_i(\cdot)$.

Solution Concept and Strategies. We are looking for an ex-ante principal-preferred subgame-perfect Nash equilibrium (SPE) of the repeated stage game.

We consider two distinct cases: The first is the *commitment case*. Here, the principal chooses her strategy at the beginning of the game and publicly commits to it. The parties then play the principal-preferred SPE, taking P 's strategy as given. The second is the *no-commitment case*. Here, the principal cannot commit a priori. Instead, we look for the principal-optimal SPE of the three-player repeated stage game.

⁵In line with the repeated games literature, we assume that the principal can use a public randomization device that—although not strictly necessary—simplifies some proofs.

A *strategy* describes a player’s full contingent plan of action. We describe it by the function $s(\cdot)$ for the principal, $y_F(\cdot)$ for party F , and $y_E(\cdot)$ for party E . We let $\sigma := (s(\cdot), y_F(\cdot), y_E(\cdot))$ describe a strategy profile, and we call a strategy profile that constitutes an SPE a *contract*, using C to denote a generic one. Our object of interest, the *optimal contract*, C^* , is the principal’s ex-ante preferred contract. The only difference between the commitment and the no-commitment case is that in the former, the principal’s strategy $s(\cdot)$ need not be incentive compatible, whereas in the latter it must be.

Preliminaries

Before moving to the analysis, we make a few preliminary observations that are useful to keep in mind.

No Transfers. Contrary to standard principal-agent models, our setting is one of “horizontal disagreement.” Thus, we implicitly assume our players have a strong mission-driven agenda, making a setting without monetary transfers salient (see, e.g., the discussions in March (1962), Tirole (1994), and Gibbons (2020)). There is one caveat to that assumption, however: If a principal holds decision-making power herself (and cares about it), we may think of b , the leadership rent, as the number of decisions she has to delegate to an agent. But then, the principal may want to increase that number (at her own expense). One could model that through direct utility transfers. Whether utility transfers benefit the principal or not depends on the parameters. A committed principal only benefits from utility transfers when she is sufficiently centrist. Without commitment, she always benefits from transfers when trying (and needing) to punish the Friend. We provide a formal discussion in Online Appendix J.

Implementable Contracts. A contract describes an equilibrium of the game. It must, therefore, be incentive compatible. Parties should prefer choosing the policy prescribed by the contract over their best deviation. If party k complies with contract C ’s recommendation, $y_k(h)$, at a history h , he will choose policy $y_k(h)$, and then expect an on-path continuation payoff $w_k(C, h)$. If, on the contrary, he deviates—e.g., by picking his preferred policy, $y_k = \theta_k$ —we switch to his penal code (Abreu, 1988). Party k ’s continuation payoff from entering his own penal code is $\underline{w}_k$.⁶ Hence, a policy y is enforceable for k within contract C at history h if and only if

$$-|y - \theta_k| + \beta w_k(C, h) \geq \beta \underline{w}_k. \quad (\text{DEC})$$

Following Levin (2003), we refer to (DEC) as k ’s *dynamic enforcement constraint*. A necessary condition for a strategy profile to be a contract is that (DEC) holds at all histories. If the principal has commitment power, the condition is also sufficient.

⁶Note that k ’s penal code is a contract, C^k , in itself. We use the shorthand $w^k \equiv w(C^k)$ to describe C^k ’s payoff vector. Each player has their own penal code and j ’s expected payoff from entering k ’s penal code is w_j^k . To emphasize that w_k^k is indeed k ’s worst continuation payoff we use $w_k^k \equiv \underline{w}_k$. Simple penal codes that only condition on the (last) deviator’s identity suffice because information is perfect.

Parties' Disagreement. Parties disagree in two dimensions: First, they disagree along the policy dimension. There, they distribute a constant utility of -1 across them. Second, they compete over the per-period leadership rent, b . Taken together, parties' utility in every period sums to $b + u_{F,t}(y) + u_{E,t}(y) \equiv b - 1$. We make two observations based on this *constant-sum property*.

Observation 1 (Pareto Efficiency). A principal-optimal contract is Pareto optimal.

The observation follows as *any* equilibrium is Pareto optimal across parties by the constant-sum property. A principal-optimal equilibrium maximizes total welfare.⁷

Our second observation concerns the role of the principal, which is grounded in her power, m , to influence the selection process. Consider a contract that replicates the Nash equilibrium of the stage game (NES hereafter) in every period. In this contract, the principal endorses her Friend, $s = -m$, and any party k chooses their bliss point $y_k = \theta_k$.

Observation 2 (No P -Power). NES is the (essentially) unique contract if $m=0$.

This is the classical min-max result that follows from the constant-sum property. If $m = 0$, the principal has no power. Thus, to achieve moderation, parties have to coordinate on concessions. But then, any policy concession by one party must be compensated through future policy concessions by the other. Because parties are impatient, future concessions must exceed current ones, which in turn require even larger concessions after—as this is impossible, parties can only agree to polarize.

Note, however, that if $m > 0$ but $b \rightarrow 0$, the principal can use endorsement to her advantage—especially in the commitment case. By threatening parties to support the opponent unconditionally, the principal can extract policy concessions even when leadership rents vanish. For instance, by moderating the Enemy, she can demand concessions from the Friend under the threat of withdrawing moderation once the Enemy returns to leadership. In the first part of the analysis, we will shut down this cross-subsidization channel and focus on an example in which the principal fully aligns with the Friend, $\theta = \theta_F = 0$. As we see, a committed principal still strictly improves upon NES for $\beta > 0$ and $m > 0$ even if $b \rightarrow 0$. The picture changes absent commitment. In the second part, we analyze the full model. We show that there, cross-subsidization interacts with the leadership rent, b , both with and without commitment.

3 Example with full alignment

We begin by analyzing the special case of a fully aligned Friend, $\theta = \theta_F = 0$.

3.1 Example: Commitment

The NES contract. A candidate for the principal would be the NES contract. This strategy maximizes the probability that the Friend takes the lead, and he then chooses the

⁷Because direct utility transfers are not possible, the reverse is not true: there may be Pareto-optimal equilibria that do not maximize total welfare and are thus not principal-optimal.

principal's bliss point, $y_F = 0$. However, the Enemy, whenever he leads, has no incentive to choose anything but his bliss point, $y_E = 1$. But is persistent NES ever optimal?

The Optimal Contract. It turns out that NES is never optimal. Instead, (almost) the opposite is true: the principal continuously embraces the Enemy once an initial attempt to exclude him has failed. More formally, the optimal contract has two phases on the equilibrium path: an initial *exclusion phase* that lasts *until the first time the Enemy leads*. Then it switches to a stationary *embracing phase*. To characterize the optimum, we specify the contract in each phase and determine how deviations are punished. With commitment and $\theta = 0$, the relevant punishment is that of a deviating Enemy, who is punished by NES.⁸

Proposition 1. *Suppose $\theta = 0$ and a principal with commitment. For any (m, β, b) , the optimal contract has two phases.*

Exclusion Phase. *The principal fully endorses the Friend, $s^0 = -m$, and the Friend chooses $y_F^* = 0$. The exclusion phase lasts until the Enemy leads for the first time.*

Embracing Phase. *From that point on, the contract is stationary. The principal fully endorses the Enemy, $s^* = m$, the Friend chooses $y_F^* = 0$, and the Enemy moderates to*

$$y_E^* = \max \left\{ 1 - \frac{2\beta m(b+1)}{1 - \beta(1/2 - m)}, 0 \right\}.$$

Proposition 1 shows that the principal switches (on-path) strategies midgame. The game begins with an exclusion phase, which lasts until nature selects E for the first time. Now, we transition irreversibly to the embracing phase. The principal asks the Enemy to moderate and promises full endorsement in all future periods in return. The Enemy complies to secure that endorsement.

The intuition behind this result comes from two observations: the constant-sum property across parties and the fact that the principal's payoff equals the Friend's payoff net of his expected leadership rents (i.e., b in periods when he leads and 0 otherwise). The principal can increase the Enemy's continuation payoff by promising him marginally more endorsement in the future. This promise gives the Enemy two benefits: First, by being in power more often, the average policy in the future is closer to his bliss point. For the sake of the argument, say the change in the expected policy increases the continuation utility of the Enemy by $\gamma > 0$. Second, the Enemy's average leadership rent in the future increases, resulting in a continuation utility gain of, say, $\omega > 0$. Thus, the Enemy's total continuation payoff increases by $\gamma + \omega$.

The principal only suffers from the Enemy's future policy choices, but not from distributing the leadership rent toward him. Her continuation utility, therefore, decreases only by γ . However, she can ask the Enemy to pay back his *total* utility gains ($\gamma + \omega$)

⁸The Friend can be left unpunished as he chooses his static optimum on-path.

through *today's policy choice*, which fully enters the principal's utility. Thus, the principal gains $\gamma + \omega$ today and loses γ in the future—a profitable trade-off.⁹

Naturally, these considerations only become relevant once the Enemy leads for the first time. Past endorsements are sunk for the leading party and, as long as there are no incentives to provide to the Enemy, the principal aims to exclude him.

3.2 Example: No Commitment

The previous results rely on the principal's promise to endorse the Enemy in the future if he moderates today. However, this promise may not be credible if the principal lacks commitment power. After all, she always prefers her Friend to lead. So, she may renege right after the Enemy has taken his action and endorse her Friend instead.

Note that this decision relates to the principal's *interim* continuation payoff, that is, *after* the Enemy chose his first policy. Although we know from Proposition 1 that more embracing of the Enemy is always better *ex ante* for the principal, she may renege on her promised endorsement if: (i) her *interim* continuation payoff is not attractive enough or (ii) the consequences of renegeing are not severe enough. As in any repeated game, the principal's dynamic enforcement constraint captures both issues: any contract C at any history h must satisfy

$$w_P(C, h) \geq \underline{w}_P, \quad (\text{DEC}') \tag{1}$$

where $\underline{w}_P$ is the principal's payoff from her penal code. Although Section 4.2 gives more details on the construction of penal codes, it is instructive to see that NES is not the harshest credible punishment for P . Recall that NES is the Enemy's worst continuation play and thus—by the constant-sum property—the Friend's best. Since Friend and principal are aligned in policies, such a penal code is too mild. Instead, to punish the principal most effectively, the Friend chooses a policy $y_F > 0$, i.e., to the right of the principal's bliss point. This is demanding for the Friend, who needs to be (at least) indifferent between continuing the principal's punishment and deviating to his bliss point $y_F = 0$, which triggers his own penal code. As a result, another complexity appears: the penal codes of Friend and principal are interconnected. We have to characterize the Friend's punishment although it is irrelevant on path. It turns out that a worst punishment for the Friend has two phases: upon deviation, the principal fully embraces the Enemy, demanding less moderation from him (the punishment phase), and, upon the Friend's return to power, continuation play switches back to the on-path embracing phase (the back-to-business phase).¹⁰

Using our construction, we find that if the principal lacks commitment power, either the optimal contract is observationally identical to the optimal commitment contract, or

⁹Policy $y_E^* < 1$ even when $b \rightarrow 0$. This is because even in that limit, the principal could trade policy for office. However, since $\omega \rightarrow 0$, all such trades would be almost 1-to-1. The equilibrium would give the principal the same payoffs as NES. Therefore, whenever $b > 0$, the principal strictly improves over NES.

¹⁰We provide further details in Section 4.2.

NES is the optimal contract.

Proposition 2. *Assume a principal without commitment, and suppose $\theta = 0$. The optimal contract is observationally equivalent to the optimal commitment contract if and only if*

$$b \geq \bar{b}_0 := \frac{(1 - \beta)^2}{\beta(2 - \beta(1 + 1/2 - m))(1/2 + m)}.$$

Otherwise, NES is the only available contract.

For $b \geq \bar{b}_0$, the following penal codes support the equilibrium:

- *the NES contract for the Enemy;*
- *a stationary contract with $y_E = 1 > y_F = \hat{y}_F > 0$ and $s = -m$ for the principal;*
- *a punishment phase with $\hat{y}_E \geq y_E^*$ and $s = m$ that lasts until F leads again; followed by a return to the on-path embracing phase for the Friend.*

For extreme values of b , the results of Proposition 2 are straightforward. If $b \rightarrow \infty$, the policy is second-order to the Enemy. He is willing to choose $y_E^* = 0$ in exchange for endorsement, so endorsing him is costless for the principal. If $b \rightarrow 0$, the game between principal and Enemy becomes almost constant-sum, leaving no room for dynamic incentives.

For intermediate values of b , however, two opposing effects are interacting. On the one hand, there is a marginal cost: the greater the endorsement for the Enemy, the lower the Friend's chances of leading and choosing $y_F^* = 0$. On the other hand, there is an inframarginal gain: the greater the endorsement promised to the Enemy, the more the principal can ask for his moderation.

Together, these opposing effects imply that the principal's *interim* continuation payoff is (strictly) convex in the level of endorsement she promises to the Enemy. When endorsement is low, the Enemy's policy moderation is small. Thus, by increasing the Enemy's chances of leading, undesired policies become more likely—the marginal cost is large. However, as endorsement rises, so does the Enemy's moderation, reducing the marginal cost. The inframarginal gain of policy moderation in exchange for endorsement is often constant in the level of endorsement. Therefore, there is no interior optimum: in the embracing phase, the principal's optimal strategy is either NES or full embracing.

Proposition 2 establishes, furthermore, that if $\theta = 0$, then *once full embracing ceases to be optimal, the only option left is NES*—in other words, the principal entirely loses her ability to make credible promises. To see why, we need to walk through the following (slightly technical) argument.

The key observation is that the Friend's payoff consists of both leadership rent and policy payoff. The latter is identical to the principal's payoff because $\theta = 0$. Moreover, full embracing of the Enemy is strictly better *ex ante* for the principal. Therefore, no embracing can only be optimal if full embracing violates the principal's (DEC').

Now, consider the knife-edge case in which full embracing makes (DEC') hold with equality. This means the principal's *interim* continuation payoff—at the embracing phase—

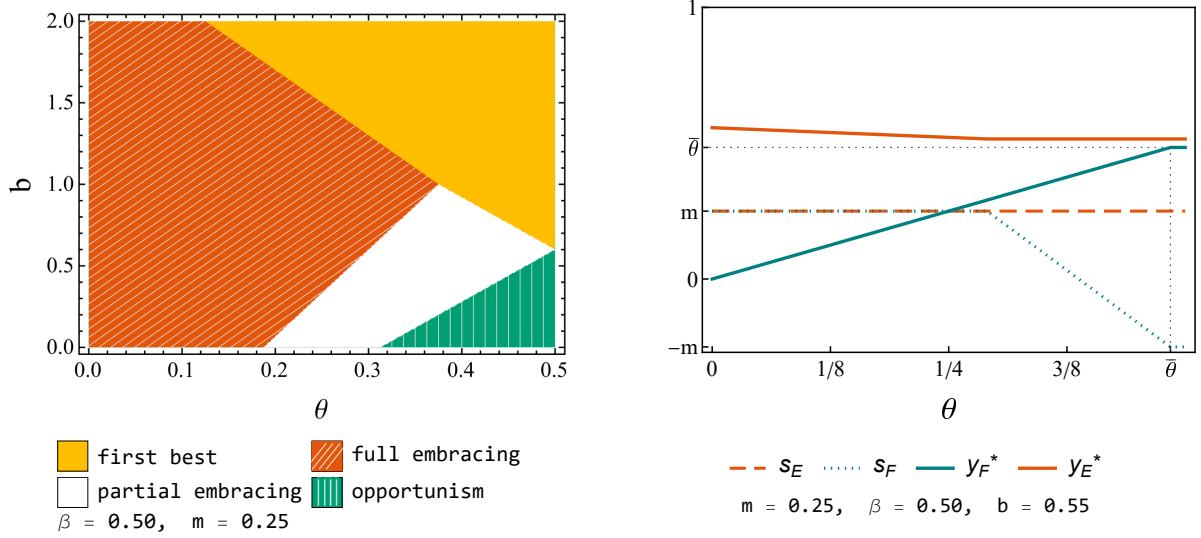


Figure 2: **Embracing Phase under Commitment.** *Left:* Regimes in the (θ, b) space. Outside the first best, *full embracing* ($s_F = s_E = m$) is optimal for low θ , *partial embracing* ($s_E = m > s_F > -m$) for intermediate θ , and *opportunism* ($s_E = m, s_F = -m$) for large θ . *Right:* Equilibrium strategies as a function of θ for given (m, β, b) . Dashed/dotted are the principal's strategies, solid those of the agents. The Friend's strategy y_F^* follows the 45-degree line until $\theta = \bar{\theta}$, then stays constant at $\bar{\theta}$.

is identical to her worst (ex-ante) payoff, $\underline{w}_P$. Consequently, the Friend's interim *policy* continuation payoff is identical and thus also $\underline{w}_P$ because his policy preferences equal those of the principal. Moreover, since the Enemy is fully embraced, the Friend's expected leadership rent in the continuation game is minimized. Hence, in the embracing phase, the Friend obtains his worst continuation payoff, that is, his penal code. Crucially, this implies that the Friend's threat of punishing the principal harder than NES (which, recall, depends on him choosing $y_F > 0$) becomes inseparable from the implementability of full embracing. Once the principal can no longer credibly promise to fully embrace the Enemy, the Friend can no longer credibly threaten the principal with policies to the right of their common bliss point. In short, the collapse of full embracing brings down both the commitment optimal contract and the principal's penal code. We revert to NES for $b < \bar{b}_0$.

Key to Propositions 1 and 2 is that the Friend needs no incentives to choose the principal's bliss point. As we will see next, once there is (sufficient) disagreement between Friend and principal, our results change.

4 Analysis

4.1 General Model: Commitment

We now turn to the general case, $\theta \in [0, 1/2)$, in which there may also be disagreement between Friend and principal. Then, the Friend would only choose the principal's bliss point if incentivized accordingly. We begin with the embracing phase. Figure 2 illustrates the findings we develop next.

Embracing. As in any dynamic setting, the principal can incentivize the Friend by (i) increasing the Friend’s on-path utility upon compliance, or (ii) threatening punishment if he deviates. In particular, a principal with commitment power can always threaten to min-max the Friend, for example, through promising unconditional and full endorsement to the Enemy, who then chooses his bliss point $y_E = 1$.

If θ is close enough to 0, the punishment threat—option (ii)—is sufficient to incentivize the Friend. The reason is intuitive: Both principal and Friend benefit from the Enemy’s moderation. Conditional on the Enemy leading, their incentives are aligned (as long as $y_E \geq \theta$). The Friend fears the Enemy choosing his bliss point 1 instead of the on-path $y_E^* < 1$. Hence, when he leads, the Friend has an incentive to choose $y_F = \theta$ even absent any promise of future endorsement. Moreover, since moderating the Friend’s choices benefits the Enemy, the principal can also ask the Enemy to moderate further. This *cross-subsidization effect* ensures that if θ is close to 0, the principal fully embraces the Enemy in the embracing phase (i.e., after the Enemy leads for the first time).

More formally, the optimal contract requires the Enemy to moderate until his (DEC) binds. The constant-sum property then implies that the losses of the Friend when the Enemy leads are bounded. In fact, the Friend attains his largest possible continuation payoff *conditional on the Enemy leading*. If θ is small, that promised utility is enough for the Friend’s (DEC) to hold with slack if he chooses $y_F = \theta$. The remaining economic forces are identical to the example. Yet, there exists a threshold value

$$\underline{\theta} := 2 \frac{\beta}{1-\beta} (b+1) m \beta \left(\frac{1}{2} + m \right),$$

such that, if the principal’s bliss point $\theta > \underline{\theta}$, cross-subsidization is no longer enough to incentivize the Friend to choose $y_F = \theta$. Then, the principal needs to provide additional, direct incentives for the Friend to make him moderate to $y_F = \theta$. The optimal way to provide such incentives is to offer a leading Friend to tilt the endorsement strategy in his favor in the *next* period. Thus, during the embracing phase, the principal selects different endorsement strategies, s_F, s_E , depending on who has led last. Not surprisingly, it remains optimal to exploit the cross-subsidization effect and require the Enemy to moderate until his (DEC) binds. Since—by the constant-sum property—the Friend is indifferent between all continuation plays that give the Enemy a certain payoff, the principal proposes the principal-optimal continuation contract that puts the Enemy at his (DEC). Using our reasoning from the example, it is thus optimal to set $s_E = m$ in that case.

We also know from the example that, absent the problem of incentivizing the Friend, the principal would like to incentivize the Enemy’s moderation by setting $s_F = m$. Thus, the only reason to lower s_F is that, otherwise, the Friend would never choose $y_F = \theta$ (i.e., the Friend’s (DEC) fails for $y_F = \theta$ and $s = m$). It is therefore optimal to set s_F such that F ’s (DEC) binds conditional on playing $y_F = \theta$. This leaves the Enemy’s incentives for moderation unchanged because, invoking the constant-sum property once more, if F ’s

(DEC) binds then E 's continuation payoff is maximal conditional on the Friend leading.

Quite intuitively, the relationship between θ and the optimal s_F^* is monotone. The larger θ , the more incentives the Friend needs to select $y_F^* = \theta$, so s_F^* declines. Depending on parameters, it may be the case that there is a second threshold

$$\bar{\theta} := \frac{1 - \beta(1/2 + m)}{\beta(1/2 + m)} \underline{\theta}$$

in the domain of θ such that s_F^* hits the boundary $s_F^* = -m$. Before turning to this case, we formalize the discussion until now and characterize the optimal contract when $\theta < \bar{\theta}$.

Proposition 3. *Suppose $\theta < \bar{\theta}$ and parameters are such that the first best is not feasible. The optimal contract has two phases.*

Exclusion Phase. *The principal fully endorses the Friend, $s^0 = -m$, and the Friend chooses $y_F^* = \theta$. The Friend's (DEC) has slack. This phase ends with the Enemy's first lead.*

Embracing Phase. *From this point on, the contract is stationary. The Friend chooses $y_F^* = \theta$, and the Enemy, $y_E^* \in (\theta, 1)$. The principal's endorsement strategy depends on θ :*

- *If $\theta \leq \underline{\theta}$, she fully endorses the Enemy, $s_E^* = s_F^* = s^* = m$.*
- *If, instead, $\theta \in (\underline{\theta}, \bar{\theta})$, she fully endorses the Enemy if he led last, $s_E^* = m$, but selects an interior endorsement $s_F^* \in (-m, m)$ if the Friend led last.*

Both s_F^ and y_E^* (weakly) decrease in θ . A leading Enemy's (DEC) always binds; a leading Friend's (DEC) binds if and only if $\theta > \underline{\theta}$ and the contract is in the embracing phase.*

Deviations are punished by unconditional full endorsement for the non-deviator and full polarization in policy choices.

The exclusion phase and its intuition are unchanged compared to the example. The Friend is willing to comply, because he receives a higher continuation value during the exclusion phase than during the embracing phase—the Enemy has no agency, so his incentives are irrelevant.

Opportunism. We now turn to $\theta \geq \bar{\theta}$. At $\theta = \bar{\theta}$, the reasoning of Proposition 3 continues to apply, but the principal has exhausted all her resources: she needs to fully endorse whoever led last to uphold that period's incentives for moderation. She promises the Friend full endorsement after leading, $s_F = -m$, and, in exchange, he moderates exactly to $y_F^* = \bar{\theta}$. By the symmetry of the principal's strategy, the Enemy chooses $y_E^* = 1 - \bar{\theta}$ in exchange for $s_E = m$. Since the principal is now *opportunistic*, the exclusion and embracing phases become almost indistinguishable. The principal follows her genuine preference for the Friend only in the very beginning, before any party acts.

For $\theta > \bar{\theta}$, the principal is incapable of incentivizing any party to choose her bliss point θ . As a result, the parameter θ has no payoff relevance to the parties and the optimal contract is identical for all $\theta \in [\bar{\theta}, \frac{1}{2})$. By the familiar logic from above, since the Enemy is at his (DEC), there is no point in decreasing $s_E^* = m$ to provide incentives to the Friend.

Hence, the principal is opportunistic. Parties moderate as much as they can be moderated, which is $y_F^* = \bar{\theta}$ and $y_E^* = 1 - \bar{\theta}$, independently of θ .

Proposition 4. *Suppose $\theta \geq \bar{\theta}$. After an initial endorsement of the Friend, $s^0 = -m$ in the first period, the contract is stationary. On-path, the principal fully endorses whoever led last, i.e., $s_F^* = -m$ and $s_E^* = m$. A leading Friend chooses $y_F^* = \bar{\theta}$ and a leading Enemy chooses $y_E^* = 1 - \bar{\theta}$. Any leading party's (DEC) binds.*

Deviations are punished by unconditional full endorsement for the non-deviator and full polarization in policy choices.

First-Best Contracts. To close the characterization, we briefly discuss first-best contracts. Because the Friend's bliss point is closer to the principal's than the Enemy's, a first-best contract is, intuitively, possible only if the principal can ensure $y_E^* = \theta$.

From our previous discussion, we can derive an immediate necessary condition for the first best: $\theta \geq 1 - \bar{\theta}$. As we know from Proposition 4, $1 - \bar{\theta}$ is the furthest we can moderate the Enemy's (average) action by putting him to his (DEC) while maximizing his chances to lead next period, and promising him the best possible continuation value if he is not selected.

However, the condition is only necessary, not sufficient. The reason is that to promise the Enemy the best continuation value conditional on not leading, the Friend may need to choose a policy $y_F > \theta$, which violates the first-best presumption.

Thus, there is a second, necessary condition: $\theta \geq 1 - \check{\theta}$, and

$$\check{\theta} := \beta \left(\frac{1}{2} + m(1 + 2b) \right),$$

where $\check{\theta}$ is the symmetric version of $\bar{\theta}$: the furthest distance the Friend and the Enemy are willing to symmetrically move in policy from their respective bliss points toward each other. Our second condition ensures that even if at $y_F = \theta$ the Friend has slack on his (DEC), the Enemy is willing to choose $y_E = \theta$. Jointly, the two conditions are sufficient to achieve first best.

Proposition 5. *The principal achieves her first best if and only if $\theta \geq \max\{1 - \bar{\theta}, 1 - \check{\theta}\}$.*

The comparative statics are intuitive. As b grows large, the Enemy becomes primarily power hungry. He is responsive to the threat of losing endorsement and willing to concede on policy to the principal; both first-best thresholds converge to $-\infty$. As m grows, the principal's endorsement becomes more valuable and punishment more severe. So, provided β and b are not too small, the first best becomes feasible. However, even in the limit $m \rightarrow 1/2$, the principal can obtain the first best only if the Friend is sufficiently close. Otherwise, even the promise of an everlasting exclusion phase cannot induce the Friend to moderate up to θ . Finally, as β increases, more outcomes become implementable by the standard logic of repeated games. Even so, note in the example of $\theta = 0$ that, as $\beta \rightarrow 1$, it is not guaranteed the first best is always implementable—we need both m and b to be sufficiently large.

Polarization. The next proposition shows that a more balanced principal is more effective at achieving moderation.

Proposition 6 (Polarization). *The distance between the parties’ equilibrium policies, $y_E^* - y_F^*$, decreases as the principal’s preferences become more balanced.*

This result highlights the moderating role of the centrist principal. It introduces a novel channel in the existing literature on repeated policy trading (see, Dixit et al., 2000; Acemoglu et al., 2011; McClellan and Liao, 2023; Invernizzi and Ting, 2024). These papers consider settings without a principal but assume curvature in parties’ utility. By contrast, our parties face a constant-sum game between them, which rules out bilateral policy trading. Their mechanism and ours are complementary: introducing utility curvature would reinforce our channel, implying that the principal’s embracing strategy generates even greater moderation.¹¹

The Symmetric Case. Before turning to the model without commitment, it is worth discussing the case of a fully balanced principal, i.e., $\theta = 1/2$. Naturally, in this case, the meaning of “Friend” and “Enemy” is lost. However, the results are (almost) identical to those of the limit case $\theta \rightarrow 1/2$. To see this, recall from Proposition 4 that the maximum moderation the principal can induce from either party is a distance $\bar{\theta}$ away from their respective bliss points. This cutoff is, importantly, independent of the principal’s bliss point θ . Now, if parameters are such that $\bar{\theta} \geq 1/2$, then a fully balanced principal can achieve first best. Otherwise, she offers the opportunistic contract described in Proposition 4, which—after the first period—is fully symmetric.

The only difference arises in this first period: while a biased principal has a *strict* incentive to endorse the Friend, a fully balanced principal is indifferent.

Proposition 7. *In the symmetric case, $\theta = 1/2$, either a first-best contract is optimal, or the opportunistic contract described in Proposition 4 is optimal. In either case, the principal is indifferent whom to endorse in the first period.*

4.2 General Model: No Commitment

We now remove the principal’s commitment power. Recall the principal’s dynamic enforcement constraint,

$$w_P(C, h) \geq \underline{w}_P, \quad (\text{DEC}')$$

which implies that, at any history h , the principal prefers to continue with contract C over switching to her worst continuation contract.

As in the example of $\theta = 0$, whether our results extend to the no-commitment case depends on two conditions: (i) the principal’s desire not to renege on a given *on-path* promise, and (ii) her ability to promise a harsh punishment if the game moves *off path* because a party deviates. On the equilibrium path, (DEC’) becomes relevant when the

¹¹In Online Appendix I, we sketch the formal argument underlying this conclusion.

contract requires the principal to endorse the Enemy; off path, when she wants to punish a deviating Friend. The reason is straightforward: the principal genuinely prefers the Friend to lead, so endorsing the Enemy strains her incentive constraint.

Punishments

How strained those incentives become depends, of course, on the ability of the two parties to punish the principal. Thus, we need to determine both the principal's and the Friend's punishment jointly. That interconnectedness of punishments often makes it hard to derive explicit, behavioral *penal codes* à la Abreu (1988) in multiplayer games.¹² In this case, however, the structure of our problem—in particular the constant-sum property between Enemy and Friend—allows us to explicitly derive each player's penal code.

Punishing the Enemy. The Enemy's grim-trigger punishment from the commitment solution corresponds to the NES contract, which poses no incentive problem for the principal.

Punishing the Friend. Punishing the Friend is considerably more difficult. A literal implementation of the commitment punishment requires parties to polarize while the principal fully endorses the Enemy. Generically, such behavior is not incentive compatible for the principal.¹³ Consequently, the Friend's optimal penal code must take a different structure. To construct it, we use the *back-to-business property*: after an initial punishment phase, play returns to the embracing phase of the principal-optimal contract.

The intuition behind this property connects directly to the logic of the embracing phase itself: When the Enemy first gains the lead, the principal moderates him by offering a desirable continuation game in exchange for maximal policy concessions. By the constant-sum property, a continuation game that benefits the Enemy is harmful to the Friend. In fact, part of what makes the promised continuation game attractive for the Enemy is that the Friend, upon leading, is held to his binding (DEC) and asked to moderate. This feature carries over to the Friend's penal code: once the Friend regains the lead after a deviation, his optimal punishment prescribes the same continuation play he faces on path. The back-to-business property serves a dual purpose. For the Friend, it ensures a low continuation payoff upon returning to lead. For the principal, it guarantees a high continuation payoff because it coincides with the optimal contract—this, in turn, incentivizes the principal during the initial punishment phase.

Now, consider the punishment phase. The principal endorses the Enemy, who moderates just enough for the principal's (DEC') to hold. We get a first insight into how the penal codes of principal and Friend are interrelated: the more the principal fears her own penal code, the greater the polarization she is willing to tolerate during the Friend's punishment phase, and hence the harsher the penalty she can credibly impose on the Friend.

¹²Formally, a player's penal code is the continuation game played after they deviated. An *optimal* penal code is the harshest such punishment.

¹³It is feasible when $\theta = 1/2$, because under symmetry, Friend and Enemy are the same to the principal.

Punishing the Principal. Finally, we determine the principal’s optimal penal code. Although its concrete structure depends on parameters, two features are always present: (i) the principal “defends” by fully endorsing the Friend—after all, she prefers the Friend at all stages and has, within the penal code, nothing to lose; (ii) the Enemy fully polarizes, choosing $y_E^P = 1$.

The only remaining object to characterize is the Friend’s action, which depends on the principal’s bliss point, θ . When θ is close enough to 0, the logic from our special case $\theta = 0$ applies. To punish the principal, the Friend chooses a policy to the right of θ . How far the Friend is willing to go depends on the Friend’s own penal code, implying a second linkage between the two penal codes. However, if we consider a large θ , the Friend would need to choose a policy that is significantly to the right. If the Friend is unwilling to choose $y_F^P \geq 2\theta$, then he instead reverts to his bliss point, $y_F^P = 0$. Hence, for sufficiently large θ , the penal code becomes the NES contract.

Because the principal’s penal code has these two regimes, her punishment value is non-monotone in her bliss point, θ . If the punishment involves $y_k^P \geq \theta$, its severity diminishes as θ approaches these policy choices. By contrast, the NES punishment becomes more severe as θ increases. She moves away from the Friend’s policy, $y_F^P = 0$, and closer to the Enemy’s policy, $y_E^P = 1$, but, because $s = -m$, the former effect dominates. We see, thus, that parties can coordinate to punish the principal harshly when she is either extreme (θ close to 0) or centrist (θ close to $1/2$). But a principal that is partially aligned (θ interior) has an attenuated punishment. This observation, of course, feeds back into the principal’s ability to promise an embracing strategy.

Optimal contract

We now (behaviorally) characterize features of the optimal non-commitment contract. To highlight the role that commitment plays in our model, we relate our findings directly to those of the commitment case.

We begin by showing that the logic on the leadership rent, b , from Proposition 2 extends to the general setting: when b is large, whether the principal can(not) commit does not matter. The same holds for *any* b if the principal’s bliss point is sufficiently close to $1/2$.

Proposition 8. *For any (β, m) , there exist thresholds $\hat{\theta}(b) < 1/2$ and $\bar{b}(\theta) > 0$ such that the optimal non-commitment contract is (observationally) identical to the optimal commitment contract if $\theta > \hat{\theta}(b)$ or $b > \bar{b}(\theta)$.*

The result hinges on the credibility of punishments: First, if the principal is balanced, she has almost no intrinsic preference between parties and can therefore credibly punish both. Second, if b is large, the parties are highly responsive to the threat of losing endorsement and thus moderate their policies on the equilibrium path, benefiting the principal. Through the back-to-business property, these benefits provide the necessary carrot that makes the principal willing to incur the cost of a severe punishment phase

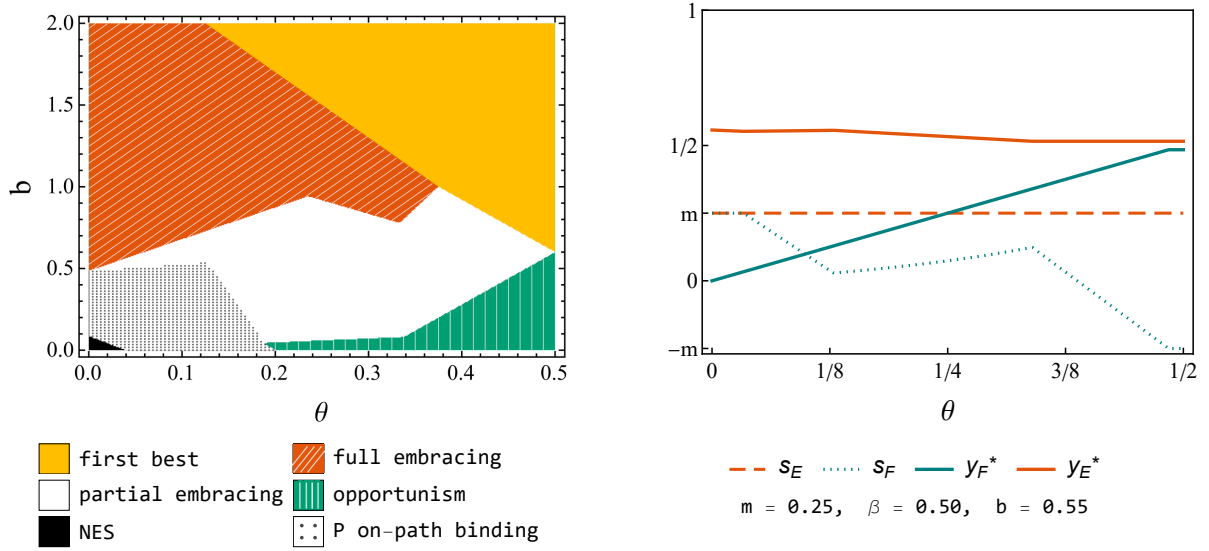


Figure 3: **Embracing Phase Without Commitment.** *Left:* Regimes in the (θ, b) space. For low (θ, b) , only the NES contract is feasible. Except in the dotted region, the binding constraint for the principal is mainly her ability to punish the Friend, not to uphold her on-path promises. Commitment is irrelevant for large (θ, b) . The full embracing region may be non-convex because the principal's punishment value is non-monotone in θ .

Right: Equilibrium strategies as a function of θ for given (m, β, b) . Relative to the commitment case, the Friend requires more endorsement, which reduces the Enemy's moderation. Kinks in s_F arise (in order) as F 's (DEC) binds on-path, P 's penal code switches to NES, and P 's commitment becomes irrelevant. The non-monotonicity in s_F reflects the non-monotonicity of the principal's punishment value.

against a deviating Friend. In turn, precisely because a harsh punishment against the Friend is credible, a harsh punishment for the principal—under which the Friend is willing to choose policies to the right of the principal’s bliss point—is also credible. Taken together, these harsh off-path punishments ensure the principal’s (DEC’) has slack on the equilibrium path, even in periods when she endorses the Enemy. Comparing the right panel of Figure 3 to that of Figure 2 illustrates that for $\theta > 3/8$, indeed the commitment solution is feasible.

This result does not tell us, however, whether the qualitative predictions of the commitment case—namely, the four *regimes* of full embracing, partial embracing, opportunism, and first best—survive for interior $\theta \in (0, 1/2)$. The next result states that for each regime, there exists a (non-knife-edge) region in which commitment plays no role. This is not a folk-theorem type of result in which direct incentives become irrelevant, and any outcome can be sustained regardless of commitment. Instead, for each regime, there is a non-trivial set of parameters for which the same incentives discussed in the commitment case are at play, and the same contract is optimal.

Proposition 9. *Fix any (β, m) . For any regime that is optimal under commitment, there exists a region of parameters $\tilde{\Theta} \times \tilde{B}$, with a non-empty interior relative to the parameter space, $\Theta \times B$, such that the optimal non-commitment contract is observationally identical to the optimal commitment contract if $(\theta, b) \in \tilde{\Theta} \times \tilde{B}$.*

The intuition behind this result comes from the analysis at the boundaries of the Θ space. At $\theta = 0$, Proposition 2 establishes that, generically, either the principal has slack on her (on- and off-path) constraints, or only the NES contract is feasible. Likewise, Proposition 8 applies close to $\theta = 1/2$, where these constraints have slack, too. Continuity of payoffs then implies that, for $b > \bar{b}_0$, around $\theta = 0$ the optimal contract close to these boundaries is unaffected by the principal’s ability to commit.

However, as we move toward the interior of the Θ space, lack of commitment becomes more consequential. Although moving in from either side can increase the principal’s payoffs under the commitment contract (see also Proposition 11), the principal’s optimal penal code also becomes less severe. This weakens the principal’s ability both (i) to punish the Friend, and (ii) to embrace the Enemy on path. Each problem jeopardizes the sustainability of the commitment contract, and both recur throughout the analysis.

Having discussed cases where the commitment contract survives entirely, a logical next step is to ask when our regimes survive qualitatively. For instance, we may wonder whether in the no-commitment case sharp cutoffs exist that partition the Θ space into different regimes, similar to $\underline{\theta}$ and $\bar{\theta}$.

It turns out that such cutoffs exist and generate dynamics resembling those of the commitment contract. Broadly, we can still identify the first-best region, as well as our three second-best regimes (full embracing, partial embracing, and opportunism). In addition, we have the NES region, where cooperation breaks down entirely. The left panel

of Figure 3 illustrates our qualitative regions.

For full embracing, the intuition is identical to that of the commitment case, but the region shrinks because the Friend’s punishment is milder. Moreover, the frontier of the full embracing region comes from the Friend’s (DEC), and not from the principal’s (DEC’). Hence, as one leaves full embracing, the principal begins to partially endorse the Friend, providing him with more incentives to moderate. We see then that, absent first best, the optimal contract transitions from full to partial embracing.

The logic of partial embracing is more subtle. The extent and the reason why the principal partially embraces the Enemy depend on which of her constraints is first-order: the off-path constraint when punishing the Friend or the on-path constraint when embracing the Enemy. If (DEC’) binds only off-path, the logic is as in the commitment case. Expected moderation by the Enemy alone is insufficient to induce the Friend to moderate up to the principal’s bliss point, so the principal supplements it with just enough endorsement to satisfy his (DEC) with equality (i.e., $s_F^* < m$ while $s_E^* = m$). This corresponds to the white region in the left panel of Figure 3. However, in the dotted region of that panel, the principal’s (DEC’) also binds on the equilibrium path. The principal is constrained when embracing the Enemy—after all, the Enemy in power is always short-term worse than the Friend in power. In this case, partial embracement occurs *after* the Enemy leads, not to accommodate the Friend’s incentive problem, but to accommodate the principal’s own. This implies that there may be situations in which both s_E^* and s_F^* are interior, so we need to broaden our definition of the partial embracing regime when thinking about the no-commitment case. Comparing the right panels of Figures 2 and 3 provides a visual illustration of how the optimal contract changes in that region.

The logic of opportunism is unchanged from that of the commitment case. The only difference is that, because the Friend’s punishment may be milder than the Enemy’s, parties may moderate asymmetrically. In the interior of the opportunism region, the principal’s (DEC’) (generally) has slack along the equilibrium path, even in periods in which she endorses the Enemy. As b decreases, and so does equilibrium moderation, the principal’s (DEC’) may eventually start binding, at which point the contract transitions to the dotted region of partial embracing. To satisfy her (DEC’), the principal can no longer credibly promise full endorsement to the Enemy after he led. Instead, she can offer at most partial endorsement, so that $s_E^* < m$ and $s_F^* = -m$. This weakens the cross-subsidization effect, putting additional strain on the principal’s incentives. As b falls further, her endorsement becomes increasingly tilted toward the Friend, until NES is her only remaining option. Because contract adjustments are marginal throughout, if the optimal contract ever reaches the NES region, it smoothly converges to it as b decreases.

While it is possible to characterize the cutoffs defining each regime, these thresholds are less informative than under commitment. One central issue is that they are sensitive to the precise structure of the optimal penal codes at each point and thus involve a set of

case distinctions. We therefore refrain from providing explicit formulas for them.¹⁴

To illustrate the additional complexities introduced by non-commitment, it is useful to consider the full embracing regime. We can show that, relative to the commitment benchmark, the full embracing region (weakly) shrinks. Moreover, for any given (m, β, θ) under which full embracing arises with commitment, there exists a threshold b_θ such that full embracing survives whenever $b \geq b_\theta$. Both results can be easily seen by comparing the left panels of Figures 2 and 3. Importantly, however, this threshold b_θ is non-monotone in θ : Full embracing may be optimal up to some θ' , partial embracing may then emerge for $\theta \in (\theta', \theta'')$, and full embracing may reappear for $\theta \in [\theta'', \underline{\theta}]$. The reason is the non-monotonicity of the principal's punishment value, which makes her punishment milder for intermediate values, in turn weakening the punishment of the Friend. As a result, there may exist an interval $\theta \in (\theta', \theta'')$ where the Friend is unwilling to moderate up to the principal's bliss point without receiving endorsement. In contrast, for $\theta < \theta'$ and $\theta > \theta''$, punishments are sufficiently harsh to discipline the Friend under full embracing.

To conclude the analysis, comparing our results with Proposition 2 gives a clear intuition: the comparative statics on b are less drastic when the principal is balanced. There are two reasons for this: First, because the Friend is often at his (DEC), the optimal contract adjusts even if the principal has no on-path commitment problem. Second, even when the principal has a commitment problem both on- and off-path, cross-subsidization may still allow for a better contract than NES.

5 Implications

In this section, we discuss the empirical predictions our model delivers for political economics. Although our model is deliberately stylized to isolate the core drivers of our mechanism, it captures key features of parliamentary coalition formation, or the behavior of interest groups. We organize the section as follows: we first address our main modeling choices. Second, we formulate three main empirical predictions, each illustrated with direct evidence from historical cases. Lastly, we turn to broader, more indirect, empirical predictions of our model.

5.1 Discussion of the Model

Community of Fate. A defining feature of our model is that players form a community of fate: They are long-lived and tied together through an organization. Policy decisions are payoff-relevant to everyone at any time, regardless of who takes them, and no one can be fully excluded. This differs from classical contracting problems, where replacement from a pool of identical players is feasible and the replaced player exits receiving an outside option. In a political system, replacement is difficult. While individual stakeholders may be replaced, the ideas behind parties, factions, or lobby organizations are there to last.

¹⁴We provide code that generates a complete symbolic characterization in the supplementary material. Online Appendix H provides some technical results resembling our discussion here.

Absent the community-of-fate assumption, our results would be less interesting. If, say, the principal could (at no cost) replace any party with a (closer) Friend, she would naturally do so. An interpretation of our discount factor is that with probability β the principal is stuck with the current parties and needs to find a way to work with them. A weaker version of replacement is to somehow eliminate the Enemy, e.g., by banning the party through a constitutional court. We will show (Proposition 12) that, often, it is not in the principal’s interest to do that.

Delegation and Limited Power. Delegation of policy choices is central to our setting. A centrist party selects coalition partners who will then shape policy from their cabinet posts, while an interest group provides electoral support to a party that then approves legislation.¹⁵

However, exogenous forces curb the authority of these power brokers. Unlike in standard delegation problems, our principal can influence the selection process but cannot fully control it. As a result, the principal needs a plan to deal with the Enemy in power. Our results show that it is in her best interest to address that issue proactively.

We believe this assumption captures a feature common to most competitive political systems. A centrist party in a multi-party system, for instance, sometimes acts as a kingmaker, determining whether the left or right bloc enters the government. Other times, the voters have given power directly to a certain bloc.¹⁶ Similarly, interest groups may support parties that ultimately lose at the polls.

5.2 The Cordon Sanitaire

A particular application of our setting is the entry of radical newcomers into parliament. Our results provide a theory of the dynamics of the so-called cordon sanitaire that mainstream parties may set up around radical newcomers by ostracizing them politically.

The Ally Principle vs. the Cordon Sanitaire. Our embracing the Enemy result stands, at first glance, in stark contradiction to the Ally Principle, which posits that a principal should hand power to her closest ally whenever possible. Considering the dynamics of our model, what we obtain is a temporary Ally Principle—the cordon sanitaire. It is optimal to collaborate with your closest ally as long as the Enemy is away from power. However, once the Enemy acquires agency, embracing him is optimal.

The Ally Principle has been challenged using arguments of face washing, power sharing, and the use of delegation as a way of committing to a policy. Within economic theory, Bendor et al. (2001) and Callander (2008) demonstrate, for example, how the Ally Principle can be invalidated if information acquisition is costly for the agents. Our mechanism differs. The principal starts to delegate power to the Enemy because her power to exclude him is limited. As a result, our theory offers a clear-cut prediction of *when* the Ally

¹⁵Delegation problems in coalitions are commonly described as *ministerial drift* (see Martin and Vanberg, 2004, for an introduction to the literature).

¹⁶In Online Appendix K we discuss a literal interpretation of the kingmaker model.

Principle holds, and *when* it will be abandoned in favor of embracing.

The Cordon Sanitaire in Reality. The sustainability of cordons sanitaires is a central question in recent European politics. Although the empirical literature is still developing, existing studies of European populist movements are consistent with two of our main findings: First, parties that were initially ostracized often become regular coalition partners once they enter government and moderate their policies (see Van Spanje and Van Der Brug (2007) and Bichay (2026)). Second, Albarello (2024) documents that moderation is greater when formerly ostracized parties are incorporated by more centrist partners.

Our theory also yields clear predictions on the dynamics followed by the embrace of the Enemy, addressing both *when* a cordon sanitaire breaks and the *extent* of the subsequent embrace. To illustrate the empirical relevance of these two additional findings, we discuss two historical cases. The first concerns the stabilization of the German Weimar Republic, showing how a cordon sanitaire ended at a particular juncture—namely, when the Enemy gained agency for the first time.

Case Study: The Stabilization of the Weimar Republic. The proportionality of the Weimar electoral system and the highly fragmented electorate gave a persistent pivotal role to the centrist parties—the Zentrum and the German People’s Party (*DVP*)—in the formation of all conceivable coalitions, either to their left with the Social Democratic Party (*SPD*), or with the German National People’s Party (*DNVP*) to their right.¹⁷ The DNVP rejected the republican constitution of 1919 and the Treaty of Versailles, and was accordingly excluded from government throughout the early 1920s. The exclusion of the DNVP meant that the centrist parties had to rely on the SPD, which belonged to a very different political tradition and, although committed to the republican constitution, pursued markedly different economic policies.

In 1924, the constitutional amendments required to ratify the Dawes Plan made the DNVP’s votes necessary. This necessity led to a bridge-building attitude from the centrist parties. As Peukert (1993) summarizes, the DVP leader Gustav Stresemann secured ratification of the Dawes Plan by “throwing out hints that there might be a place for the DNVP in the governing coalition.” After the plan was ratified, Stresemann

was conciliatory towards the DNVP, arguing that having been brought to accept realities, it should be allowed a share in government [...] He argued that bringing the DNVP into government would be an act of statesmanship. No party, he said, had more reason than they to dislike the DNVP: “But one cannot make domestic policy with sentimentality.” (Wright, 2002, 296ff.)

The ratification of the Dawes Plan marked the stabilization of the Weimar Republic, a period of DNVP moderation. The DNVP joined the government and, to maintain

¹⁷Our model abstracts from the fact that voters react to the parties’ policy choices, thus affecting the principal’s power in the future. During the 1920s, the centrist parties retained their pivotal role and faced a relatively stable menu of coalition options despite high electoral volatility.

its position, implicitly recognized the Weimar constitution, renewing the “law on the protection of the Republic—originally passed as an emergency measure against ‘the Enemy on the right’” (Maier, 2015). The DNVP leaders understood that participation in government allowed them to deliver benefits to their constituencies, particularly through protective tariffs for farmers. Until the collapse of the Weimar system, several power swings occurred in the Reichstag with the centrist parties alternating their collaboration with both the SPD and the DNVP. The Luther I cabinet (1925) and Marx IV cabinet (1927–28) included the DNVP, while the Marx III cabinet (1926–27) relied on SPD toleration, and the M  jller II cabinet (1928–1930) included the SPD.¹⁸

The stabilization effectively ended largely due to forces exogenous to our setting. The Great Depression fueled the rise of the Nazi and Communist parties, leaving none of the above coalition options able to command a majority. Shortly before that, the DNVP grassroots radicalized, culminating in the ousting of the leadership that had presided over the moderation process. Within a year, Stresemann died.

The extent of the embrace. While our prediction for *when* the cordon sanitaire breaks down is independent of the parameterization, behavior within the embracing phase is not. To connect our qualitative results to observable behavior, we focus on the cleanest divide of our theory: whether the principal fully endorses the Enemy during the embracing phase or not. The next proposition characterizes how the cutoff θ varies with the parameters under commitment.

Proposition 10. *The full-embracing region weakly expands in b , and β . With commitment, it also expands in m .*

Therefore, conditional on the Enemy having gained agency at least once, full-embracing behavior should be more likely in settings where office-holding is more valuable (b), the political environment supports longer relational horizons (β), and, at least for the commitment case, where the principal is stronger (m).

To relate these results to real-world settings, we briefly introduce a second historical example, drawn from postwar Italy. As we elaborate below, this example reflects a more stable environment, with a principal who is a stronger power broker and substantially closer to the Friend.

The *Apertura a Sinistra*. In postwar Italy, the Christian Democratic Party (*DC*) governed as the dominant centrist party alongside two smaller allies, the Republicans (*PRI*) and Social Democrats (*PSDI*). This centrist bloc faced a choice between two potential coalition partners: the Liberal Party (*PLI*) to its right and the Socialist Party (*PSI*) to its left. The ideological distance between the DC and the PSI was especially pronounced in the first postwar decade, when the Socialist leader Pietro Nenni maintained a unity-of-action pact with the Communists. Accordingly, throughout the 1950s, the DC excluded the PSI

¹⁸Table 1 in Online Appendix M tracks the principal’s options. It shows that the principal indeed exercised agency as both options were arithmetically feasible at times.

from government, forming a center-right coalition that commanded razor-thin majorities. When even this formula proved insufficient, DC minority governments relied on *de facto* support from neo-fascist deputies of the Italian Social Movement (MSI)—a dependence that reached its crisis point with the Tambroni government. Tambroni’s reliance on MSI votes triggered nationwide anti-fascist protests and violence, prompting the DC leadership to terminate the government. As one party leader put it, the party “no longer had the choice” of allying rightward (Leonardi and Wertman, 1989).¹⁹

The opening to the left proceeded in several steps. In 1960, the PSI abstained on a DC’s minority government led by Amintore Fanfani, and in 1962, the PSI provided external support. In December 1963, Aldo Moro formed the first center-left cabinet with PSI ministers. The DC leadership understood the coalition shift in a way consistent with our embracing strategy. They intended to attract the Socialists with “the trappings of office and the chance of carrying through some reforms,” while retaining the threat of expulsion should the PSI refuse to moderate (Ginsborg, 1990). Nenni made significant concessions on NATO, and their initially ambitious reform program was progressively diluted. As the Christian Democrat Carlo Donat-Cattin observed, the DC leadership had turned the center-left governments into

a ‘power arrangement’ in which the Socialists had been substituted for the Liberals, but in substance nothing had changed in the way that power was managed. (Leonardi and Wertman, 1989)

The Socialists became the regular coalition partner, and the Liberals did not return to government for almost a decade. This pattern is especially striking because parliamentary arithmetic made both formulas viable throughout the 1960s and early 1970s. Yet the DC consistently chose the center-left.²⁰

Comparing the Two Historical Examples. Both examples exhibit embracing and moderation, but differ in the extent of the embrace. In the Weimar Republic, the centrist bloc alternated between the DNVP and the SPD within the space of three years. The Italian Christian Democrats, however, governed with the PSI for almost a decade before briefly restoring a coalition with the PLI. These contrasting patterns are consistent with our comparative statics from Proposition 10. The first, and most important, distinction concerns the principal’s relative ideological distance from the Friend. In Weimar Germany, the SPD was also far from the centrist bloc, so the centrists needed to keep the incentives for moderation on both sides of the spectrum. In Italy, in contrast, the Liberals were significantly closer to the DC than the PSI. Second, the stability of the political environment and the influence of the principal differed. The Italian Republic, upheld by the American

¹⁹External changes also facilitated a leftward turn: the Kennedy administration reversed its opposition to Socialist participation, and Pope John XXIII adopted a stance of non-interference in coalition politics.

²⁰Tables 2-4 in Online Appendix M track the principal’s options.

international order and anchored in the DC's solid electoral base, offered a longer effective time horizon and a more influential principal than the fragile Weimar system, which was buffeted by reparations, economic instability, and rapid partisan realignment. Third, the value of holding office was markedly different. Postwar Italian governments presided over the *trente glorieuses*—a period of rapid economic growth in which office carried substantial rents, including widespread opportunities for patronage. In Weimar Germany, by contrast, office meant administering deeply unpopular austerity and reparations.

5.3 Broader Implications

We now address the broader implications of our model. We discuss how to identify relevant parameters or how, if the parameters are known, they map into economic predictions about behavior.

Principal's Welfare. We show that a principal's payoff is highest when she exhibits some bias toward a party, but the bias is not extreme.

Proposition 11 (Principal's Payoff). *Consider (β, m, b) such that the first best is unattainable for any θ . Then the principal's ex-ante payoff is non-monotone in θ with an interior maximum $\theta^{\max} \in (0, 1/2)$. If the principal has commitment power that maximum is unique and at $\theta^{\max} = \bar{\theta} < \frac{1}{2}$.*

The intuition for this result can be most easily seen in the commitment case. We start with an extreme principal, $\theta = 0$, and gradually make her more balanced. For $\theta \leq \bar{\theta}$, the Friend keeps choosing $y_F^* = \theta$. At the same time, the Enemy both gets mechanically closer to the principal and—through the cross-subsidization effect—has stronger incentives to moderate. However, for $\theta > \bar{\theta}$, the optimal contract is opportunistic (Proposition 4) and does not change as θ increases. Then moving the principal's bliss point toward the center makes her worse off when F leads and better off when E leads. However, because of the exclusion phase, alignment with the Friend is more valuable for the principal—so her ex-ante payoff decreases in θ . Removing the principal's ability to commit does not change that intuition. The only difference is that single-peakedness is not guaranteed because the principal's punishment payoffs are non-monotone in θ .

The Value of the Enemy. We now ask whether the principal benefits from the existence of the Enemy. Absent the Enemy, the principal loses the ability to moderate her Friend.

To state our result formally, first let us define

$$\Upsilon(m, b) := \left\{ (\beta, \theta) \text{ such that } \theta < \max \left\{ 1 - \bar{\theta}, 1 - \check{\theta} \right\} \right\},$$

which is the set of (β, θ) for a given (m, b) such that the first best is not attainable (see Proposition 5). We say that *the Enemy is valuable* if, ex ante, the principal prefers our baseline game to a benchmark without E , in which F has full agency in every period.

Proposition 12 (The Value of the Enemy). *Fix the parameters (m, b) such that $\Upsilon(m, b)$ is non-empty. The following holds:*

- For any β , there exists a $\tilde{\theta}$ such that $(\beta, \tilde{\theta}) \in \Upsilon(m, b)$ and the Enemy is valuable for the principal if $\theta \geq \tilde{\theta}$.
- If $2b \geq \frac{1-2m}{2m}$, then for any $\theta > 0$, there exists a $\tilde{\beta}$ such that $(\tilde{\beta}, \theta) \in \Upsilon(m, b)$ and the Enemy is valuable for the principal if $\beta \geq \tilde{\beta}$.

The Enemy’s presence allows the principal to discipline both parties. The cost is, naturally, that the Enemy occasionally chooses the policy. Intuitively, the more balanced the principal, the lower that cost, and the greater her need to discipline the Friend. If players are patient enough, the principal’s ability to moderate is strong. Therefore, she prefers a competitive setting, even when she is almost aligned with the Friend.

Proposition 12 offers a new perspective on why a polity may endure as a “community of fate”, where factions fundamentally disagree over policy directions, yet none is permanently excluded. Conventional explanations are that exclusion is too costly—e.g., because it requires violence—or precluded by institutional constraints. Our result suggests an additional mechanism: the principal may preserve political rivalry as it induces both sides to moderate in equilibrium.

Removing commitment decreases the value of the Enemy because the Friend cannot be moderated as effectively. Yet, the sufficient condition in Proposition 12 remains valid.

6 Conclusion

We provide a theory of why a weak power broker may want to embrace her Enemy: Moderating the Enemy’s policy choices is better than trying to ensure her Friend is selected more often. Moreover, the principal gains from having an Enemy that challenges her Friend, even if that implies some periods of worse policy decisions for her. The principal’s bias against the Enemy affects both her benefits from these relationships and the role commitment plays in them.

There are, naturally, several aspects of reality our model ignores. Among these are evolving preferences, the monitoring of the principal’s endorsement, and the potential role of asymmetric information. We leave those to future research.

Appendix

Notation: Histories. A history h is a vector describing past actions and realizations. The initial node is given by $h_0 \equiv \emptyset$. We say $h \subseteq h'$ if history h' is reached in a continuation game of h . If, in addition, $h \neq h'$, then $h \subset h'$. The set of histories h in which party k has been the only one leading is $\mathcal{H}_k$. The set of histories h where k leads for the first time at h is $\mathcal{H}_k^*$. To ensure readability, we abuse notation and suppress within stage game histories. For example, we write $y_k(h)$ meaning party k ’s policy when selected in the stage game starting at history h following “on-path play” by the principal in her choice $s(h)$. This

allows us, for example, to compare the mutually exclusive $y_F(h)$ and $y_E(h)$ in a compact manner. When it is clear from the context what we mean, we drop the history argument for a more compact formulation describing on-path stationary (continuation) strategies.

A Useful Lemmas

Lemma 1. *Take any contract with ex-ante value $w_P(h_0)$ to the principal. Then there exists a contract with value at least $w_P(h_0)$ such that $y_F(\cdot) \leq y_E(\cdot)$.*

Proof. Consider a contract C where at some history h , $y_F(h) > y_E(h)$. At this history, the principal's selection decision is $p(s(h))$, and the principal's continuation value before selection is $w_P(h)$. Suppose now that, in addition, either $0 < y_F(h) < \theta$ or $1 > y_E(h) > \theta$. Consider a contract $\tilde{C}$ that is identical to C except that at history h it specifies

$$\tilde{y}_F(h) = y_F(h) - \varepsilon \text{ and } \tilde{y}_E(h) = y_E(h) + \frac{1 - p(s(h))}{p(s(h))} \varepsilon$$

for some $\varepsilon > 0$. Because $|\tilde{y}_k - \theta_k| \leq |y_k - \theta_k|$, $\tilde{C}$ is a contract. Moreover, at history h the principal's value is identical because

$$\begin{aligned} \tilde{w}_P(h) - w_P(h) &= -\left(p(s(h))(|\tilde{y}_E(h) - \theta| - |y_E(h) - \theta|) + (1 - p(s(h)))(|\tilde{y}_F(h) - \theta| - |y_F(h) - \theta|)\right) \\ &= -\left((1 - p(s(h)))\varepsilon - (1 - p(s(h)))\varepsilon\right) = 0. \end{aligned}$$

because the net-present value to the principal of the game *after* completing this period's stage game is constant by construction.²¹ Now consider $y_E(h) \leq \theta \leq y_F(h)$. As before, take a $\tilde{C}$ that is identical to C but with

$$\tilde{y}_F(h) = \max\left\{\theta - \frac{p(s(h))}{1 - p(s(h))}(\theta - y_E(h)), 0\right\} \text{ and } \tilde{y}_E(h) = \min\left\{\theta + \frac{1 - p(s(h))}{p(s(h))}(y_F(h) - \theta), 1\right\}.$$

Thus, $\tilde{C}$ is a contract with $\tilde{y}_F(h) \leq \tilde{y}_E(h)$, and P 's value (weakly) increases because

$$-p(s(h))(\tilde{y}_E(h) - \theta) - (1 - p(s(h)))(\theta - \tilde{y}_F(h)) \geq -p(s(h))(\theta - y_E(h)) - (1 - p(s(h)))(y_F(h) - \theta)$$

which implies $\tilde{w}_P(h) \geq w_P(h)$. Unless either policy is at its extreme, $\tilde{w}_P(h) = w_P(h)$. $\square$

Lemma 2. *The optimal on-path policies satisfy $y_E(h) \geq \theta \geq y_F(h) \forall h$. It is without loss of generality to focus on contracts in which, if a party gets to choose y on path, (i) either he chooses the principal's bliss point $y_k = \theta$, or (ii) (DEC) binds.*

Proof. First Statement: $y_E \geq \theta \geq y_F$. Fix a contract C such that $y_E(h) < \theta$ for some on-path history h . We will augment the contract in (at most) two steps to construct an alternative that the principal prefers. First, consider an identical contract, apart from increasing $y_E(h)$ marginally. That contract improves the principal's payoff at node h by dy_E , relaxes E 's (DEC) in all histories $h' \subseteq h$, and leaves them unchanged in all periods $h'' \supset h$. It is thus incentive compatible for E . Now consider the last history $\hat{h} \subset h$ in the

²¹The "present" here is node h .

sequence leading up to h in which F leads, if such a history exists. The increase of $y_E(h)$ has tightened F 's (DEC) by $\delta(h|\hat{h})dy_E$, where $\delta(h|\hat{h})$ is the discounted probability of reaching h from $\hat{h}$.²² Thus, augmenting the contract once again by reducing the prescribed $y_F(\hat{h})$ by at most $\delta(h|\hat{h})dy_E$ restores (DEC) of F without violating E 's (DEC) anywhere. Note that if $y_F(\hat{h}) > 0$, the perturbation is always possible; if $y_F(\hat{h}) = 0$, we take the next “earlier history” at which $y_F > 0$ (note that if there is any issue with F 's (DEC) binding, such a history must exist). Now, observe that the first and second augmentation jointly are at worst ex-ante payoff neutral to P because at node $\hat{h}$ her gains and losses equate. Hence, assuming $y_E(h) \geq \theta$ is without loss. The case for $\theta \geq y_F$ is analogous.

Second Statement: Either (DEC) binds or $y_k = \theta$.

Definition 1. History h with property Π is a *first history with property Π* if no history $h' \subset h$ has property Π . It is a *last history with property Π* if no $h' \supset h$ has property Π .

Step 1. k 's dynamic enforcement constraint upon first selection. Take an arbitrary contract C with $y_E(\cdot) \geq y_F(\cdot)$ and assume there is a first history in which k leads, $h \in \mathcal{H}_k^*$, and k 's (DEC) has slack. Then consider a contract $\tilde{C}$ identical to C apart from $\tilde{y}_k(h)$ which is chosen such that $|\tilde{y}_k(h) - \theta| < |y_k(h) - \theta|$ and k 's (DEC) at h is not violated. Since $\tilde{y}_k$ implies greater moderation, $\tilde{C}$ is strictly preferred by P and $-k$. Now, either $\tilde{C}$ exists for $\tilde{y}_k(h) = \theta$, or there is a $\tilde{y}_k(h) \neq \theta$ such that k 's (DEC) binds. We iterate this procedure for any $h \in \mathcal{H}_k^*$ and obtain that an optimal contract exists in the class $\mathcal{C}^1$ in which the second statement holds for any k and $h \in \mathcal{H}_K^*$.

Step 2. k 's dynamic enforcement constraint upon re-selection. Consider a contract C^1 in the class of $\mathcal{C}^1$ and assume that there is at least one history h such that k 's (DEC) has slack. Take the first on-path history h with that property.

Consider then a candidate contract C^2 which is identical to C^1 with the exception that at history h , we pick the associated $\tilde{y}_k(h)$ such that either $\tilde{y}_k(h) = \theta$ or k 's (DEC) binds. Using the same arguments as in step 1, C^2 is incentive compatible for P and $-k$. However, it may not be incentive-compatible for k at a previous selection.

To construct a contract incentive compatible for k , take the last on-path history $h' \subset h$ at which k led before h . Consider now a contract C^3 identical to C^2 , but with the difference that at history h' , party k chooses $\tilde{y}_k(h')$ such that $|\tilde{y}_k(h') - y_k(h')| = \delta(h|h')|\tilde{y}_k(h) - y_k(h)|$ and $|\tilde{y}_k(h') - \theta_k| < |y_k(h') - \theta_k|$.²³ Such a $\tilde{y}_k(h')$ exists by construction and implies that k decreases his moderation at history h' by exactly the net-present value of his own increase in moderation at history h . Thus, k is indifferent at history h' between contracts C^3 and C^1 and, by construction, k 's (DEC) at node $h' \subset h$ holds in contract C^3 . Further, since these concessions are equivalent in history- h' net-present values, they are also equivalent

²²Formally, if h is a period- t history and h' a period- t' history, then $\delta(h|h') := \beta^{t-t'} Pr(h|h')$.

²³Recall the definition of $\delta(\cdot|\cdot)$ from Footnote 22. As in the proof of the first statement, if $|\tilde{y}(h') - \theta|$ cannot be chosen because $y_k(h')$ is too close to θ , we can move to an earlier history h'' with the same argument and move $\tilde{y}(h'')$ at that history.

at any $h'' \subset h'$. Two implications follow from this feature: first, C^3 is ex-ante payoff equivalent to C^1 for P and second, C^3 is incentive compatible if the initial C^1 is.

Applying Step 2 iteratively at any history in which k leads implies the statement. $\square$

We now derive the ex-ante worst payoff for a party.

Definition 2 (Grim Trigger punishment). The contract of party k is said to be *grim trigger* if it prescribes that the principal fully endorses $j \neq k$ in any continuation game and both parties choose their respective bliss points whenever leading, $y_F = 0$ and $y_E = 1$.

Lemma 3. *In a contract with commitment, a party's worst ex-ante continuation payoff is attained by grim-trigger punishment and equals*

$$\alpha := \frac{(1/2 - m)b - (1/2 + m)}{1 - \beta}.$$

Proof. Let $\underline{w}_k$ be party k 's worst ex-ante continuation payoff. Under grim-trigger punishment, the principal and the other party ensure that the probability of achieving the lower bound of the stage payoff, -1 , is maximized. Moreover, in a punishment contract, when leading, the punished party's payoff is at least $b + \beta \underline{w}_k$. Thus, the lowest ex-ante payoff must satisfy

$$\underline{w}_k \geq (1/2 - m)(b + \beta \underline{w}_k) + (1/2 + m)(-1 + \beta \underline{w}_k) \Leftrightarrow \underline{w}_k \geq \frac{(1/2 - m)b - (1/2 + m)}{1 - \beta} =: \alpha.$$

Grim trigger implies a payoff of α and therefore $\underline{w}_k = \alpha$. $\square$

Lemma 4. *In any optimal contract with commitment in which for some h , $y_E(h) \neq \theta$, the principal either fully endorses E at any history $\eta \supseteq h$ or F 's (DEC) binds at some $\eta \supseteq h$. In any optimal contract without commitment, the same holds unless P 's (DEC') binds.*

Proof. Consider some optimal contract C and consider a first history $\hat{h}$ of that contract at which E 's (DEC) holds with equality and $y_E(\hat{h}) \neq \theta$. Now assume the contract specifies $s(\check{h}) \neq m$ for some history $\check{h} \supset \hat{h}$.

We first show that unless F 's (DEC) binds at some $\eta \supset \hat{h}$, the principal has a strict incentive to increase $s(\check{h})$ contradicting optimality. This proves the commitment case.

To see this, observe that a *ceteris paribus* increase in $s(\check{h})$ increases E 's continuation payoff at $\check{h}$ by a *policy* term (a change in the average policy $-1/(1 - \beta)$ allocated to E moving forward from $\check{h}$) and an *leadership* term (a change in the fraction of $b/(1 - \beta)$ allocated to E moving forward from $\check{h}$). That means, the previous payoff $w_E(\check{h})$ increases to $w_E(\check{h}) + \gamma + \omega$ where $\gamma > 0$ is the policy term and $\omega > 0$ the leadership term. Moreover, by (2) and the fact that P therefore always prefers a leading F , P 's continuation payoff at $\check{h}$ decreases at most by E 's policy term to $w_P(\check{h}) - \gamma$.

The realization of $\check{h}$ is uncertain, and parties discount the future by β . However, both parties face the same uncertainty. Thus, from the perspective of $\hat{h}$, both discount changes in the continuation payoffs at a future history $\check{h}$ in the same way from their history $\hat{h}$

perspective, implying history- $\hat{h}$ continuation payoffs of $v_E(\hat{h}; E) + \hat{\gamma} + \hat{\omega}$ and $v_P(\hat{h}; E) - \hat{\gamma}$, respectively. If, in addition, $\check{h}$ occurs with positive probability and in finite time given $\hat{h}$, both $\hat{\gamma}$ and $\hat{\omega}$ are positive.

Being promised a strictly higher continuation payoff, E 's (DEC) no longer binds. He is thus willing to decrease $y_E(\hat{h})$ by $\hat{\gamma} + \hat{\omega}$ —a pure policy effect translating to an increase in P 's payoff to $v_P(\hat{h}; E) + \hat{\omega}$. Thus, we have constructed a contract that keeps E 's incentives in place, but is strictly preferred by P at any history $h \subseteq \hat{h}$. An immediate contradiction to optimality of C unless our construction violates F 's incentives.

To see the result for the no-commitment case, observe that the only difference is that an increase in $s(\check{h})$ may now violate P 's (DEC'). Unless it does, the argument above applies. That is, without commitment there cannot be any history $\check{h}$ such that $s(\check{h}) < m$ unless either F 's (DEC) or P 's (DEC') binds along the path from $\hat{h}$ to $\check{h}$. $\square$

B Example with Full Alignment

B.1 Proof of Proposition 1

Proof. Since $\theta = 0$, in any optimal contract $y_F = \theta$ with F 's (DEC) trivially satisfied. Before E leads for the first time, i.e., for $h \in \mathcal{H}_F$, P 's endorsement has no effect on E 's incentives, so $s(h) = -m$. After the first time E leads, i.e., for $h \notin \mathcal{H}_F$, by Lemma 2, either E chooses P 's preferred policy θ or E 's (DEC) binds. Applying Lemma 4, this implies $s(h) = m$. Solving for $y_E(h)$ using inequality (DEC) yields y_E^* . $\square$

B.2 Proof of Proposition 2

Proof. When the commitment contract yields the first best, the result holds trivially. For the remainder, we prove the non-first-best case.

First Statement Consider the following contract:

On the equilibrium path. The contract is identical to the optimal commitment contract described in Proposition 1.

If E deviated. The continuation contract is the NES outcome, $(y_F = 0, y_E = 1, s = -m)$.

If P deviated. The continuation contract is stationary and given by $(y_F = \hat{y}_F, y_E = 1, s = -m)$; $\hat{y}_F > 0$ and makes F 's (DEC) bind.

If F deviated. The continuation contract has two phases: (i) punishment phase and (ii) back-to-business phase. The punishment phase lasts until the Friend is selected for the first time (after his initial deviation from the on-path contract). During the punishment phase, every period's action profile is $(y_E = \hat{y}_E, s = m)$, where $\hat{y}_E$ makes P 's (DEC') bind. Once we enter the back-to-business phase, the contract is identical to the on-path contract's embracing phase.

The optimality of these penal codes is proved for the Enemy and principal by Lemma 3 and Lemma 15, respectively, and for the Friend by Lemma 17 if the optimal on-path

contract features full embracing. By construction, P is indifferent in the punishment phase of F 's penal code between choosing $s = m$ and deviating to $s = -m$. Thus, we have that

$$-p(m)\hat{y}_E + \beta \underline{w}_P = \underline{w}_P = -\frac{p(m)\hat{y}_F + (1 - p(m))}{1 - \beta}. \quad (1)$$

The first equality comes from P 's indifference in F 's punishment, the second by the definition of P 's punishment payoff. Likewise, F is indifferent in P 's penal code between choosing $y_F = \hat{y}_F$ or deviating to $y_F = 0$ when leading. That is,

$$-\hat{y}_F + \beta w_F^P = \beta \underline{w}_F. \quad (2)$$

Next, observe that because all actions are taken (because $\theta = 0$) to the right of the principal's bliss point, the Friend's continuation payoff at any history is equal to that of the principal plus the Friend's expected leadership rent. In particular, this implies

$$\underline{w}_F = \underline{w}_P + \frac{1 - p(m)}{1 - \beta} b = \frac{-p(m)\hat{y}_E + (1 - p(m))b}{1 - \beta} \quad (3)$$

with the last equality coming from substituting for $\underline{w}_P$ via (1). This yields four linear equations: two from (1) and those from (2) and (3). If these equations have a solution, it is unique and yields, in particular,

$$\hat{y}_E = 1 - \frac{(2p(m) - 1)(1 - \beta(1 + p(m)b))}{p(m)(1 - \beta)}.$$

Now, recall that in the exclusion phase of the optimal contract, the principal has no commitment problem because she supports the Friend anyhow. In the embracing phase, she aims to implement a contract that is identical to F 's penal code, apart from E selecting y_E^* instead of $\hat{y}_E$ in the punishment phase. We can find $\hat{y}_E \geq y_E^*$ if and only if

$$b \geq \frac{(1 - \beta)^2}{\beta p(m)(2 - \beta(2 - p(m)))} =: \bar{b}_0.$$

Second statement. First, by Lemma 2, we see that F 's (DEC) never binds on the equilibrium path (at best it holds with equality). This implies, invoking Lemma 4, that either $s = m$ or P 's (DEC') binds at some history on the equilibrium path. Second, since F needs no incentives to select his bliss point, we have that $y_F^* = 0$. Third, by optimality, E 's (DEC) binds at least in his first action node in the embracing phase.

We now want to show that it is without loss of generality to assume that E 's (DEC) binds at all of E 's decision nodes. To do that, take an optimal contract at which E 's (DEC) binds in his first action node, $h_t \in \mathcal{H}_E^*$, but has slack at history h_{t+1} with $h_{t+1} \supset h_t$. Note that there exists an alternative strategy profile that is equivalent in everything except that it lowers $y_E(h_{t+1})$ and raises $y_E(h_t)$ in a way such that (i) E 's (DEC) binds both at h_t and h_{t+1} ; and (ii) P 's payoff before h_t is unchanged. This strategy profile is a contract because, by reshuffling E 's policies, we have weakly improved P 's payoff at $w_P(h_t)$, implying that P 's (DEC') holds. Since F needs no incentives to choose $y_F^* = 0$, the same argument generalizes to the case in which F is selected for some number of

periods before E is selected again. Reiterating this same argument across every node in which E acts, we obtain an optimal contract such that E 's (DEC) binds at every history, which is at least as good as the optimal contract we started out with.

Consider the optimal contract we have constructed. Once the embracing phase starts, this contract is stationary such that it 'restarts the embracing phase' at every history in which E led. This implies that if P 's (DEC') binds at some on-path history, then P 's (DEC') binds at every on-path history. We can write P 's payoff for $h \in \mathcal{H}_E^*$ as:

$$v_P^*(h \in \mathcal{H}_E^*) = \theta - y_E^*(h \in \mathcal{H}_E^*) + \beta \underline{w}_P.$$

Because the contract is optimal for the principal, there exists no $y_E < y_E^*(h \in \mathcal{H}_E^*)$ that is incentive compatible for E . Thus, E is at his (DEC) which implies

$$v_E^*(h \in \mathcal{H}_E^*) = b + y_E^*(h \in \mathcal{H}_E^*) - 1 + \beta w_E^*(h \in \mathcal{H}_E^*) = b + \beta \underline{w}_E.$$

Then, optimality for the principal guarantees that there is no $w_E(E) > w_E^*(h \in \mathcal{H}_E^*)$ that the principal can credibly offer E . If there were, she would do it in exchange for lowering $y_E^*(h \in \mathcal{H}_E^*)$. Hence, the continuation play of the optimal contract attains the maximum value for E and, by the constant-sum property, the minimum value for F . Thus, an optimal penal code for F consists of restarting the embracing phase of the optimal contract. We use this punishment to construct P 's penal code following Lemma 15, which has the following condition:

$$-\hat{y}_F + \beta w_F^P = \beta \underline{w}_F = \beta w_F^*(h \in \mathcal{H}_E^*). \quad (4)$$

Similarly, P 's (DEC') binding in the embracing phase of the optimal contract implies:

$$\underline{w}_P = -\frac{(1 - p(-m))\hat{y}_F + p(-m)}{1 - \beta} = w_P^*(h \in \mathcal{H}_E^*). \quad (5)$$

And, because $y_F^* = 0$ requires no incentives and P 's (DEC') binds (at least) at the first node of the embracing phase, it is without loss to consider stationary choices in that phase, $(s^*, y_E^*(s^*))$,²⁴ where the latter is chosen such that E 's (DEC) binds. Thereafter, we refer to $w_k^*(E)$ as k 's continuation value in the embracing phase. That is,

$$-(1 - y_E^*(s^*)) + \beta w_E^*(E) = \beta \underline{w}_E = \beta \alpha,$$

where

$$w_E^*(E) = \frac{p(s^*)(b - (1 - y_E^*(s^*))) - (1 - p(s^*))}{1 - \beta}.$$

Since $y_k \geq \theta = 0$, we can express F 's payoffs as P 's payoffs plus his expected leadership

²⁴Recall from our description in the main text that (with some abuse of notation) s^* denotes the principal's strategy during the embracing phase.

rents:²⁵

$$w_F^P = \frac{1 - p(-m)}{1 - \beta}b + \underline{w}_P \quad \text{and} \quad w_F^*(E) = \frac{1 - p(s^*)}{1 - \beta}b + \underline{w}_P,$$

respectively. Plugging into (4) and (5) and solving for $\hat{y}_F$ and s^* , we obtain two solutions.

One solution is $s^* = -m$ leading to $y_E^*(s^*) = 1$ and $\hat{y}_F = 0$. This implies repetition of the NES outcome in the embracing phase of the optimal contract (and *all* penal codes). The other solution is given by $\hat{y}_F > 0$ and

$$s^* = \frac{\beta^2(b(2m+3)+4) - 4\beta(bm+b+2)+4}{2b\beta(\beta(2m-1)+2)},$$

which is monotone in b and leads to a $s^* \leq m$ only if $b \geq \bar{b}_0$. Thus, below $b < \bar{b}_0$, the only contract in which P 's (DEC') holds with equality is the NES contract.²⁶ $\square$

C The Optimal Commitment Contract

C.1 Proof of Propositions 3 to 5

Proof. We provide the full closed-form characterization of the optimal contract in Online Appendix G. We state the lemmas leading to these results here. Jointly, they prove Propositions 3 to 5.

Specifically, Lemmas 5 to 7 prove the if part of Proposition 5. Lemmas 8 to 11 prove Proposition 3. Lemma 12 proves Proposition 4 and the only-if part of Proposition 5. $\square$

Many of the lemmas work the following way: We first impose a particular contract candidate (Step 1), then we check the parties' (DEC) to ensure it is indeed a contract (Step 2), and (if needed) we address the exclusion phase (Step 3). If the proof of a Lemma deviates from this structure, we comment on it here explicitly.

We make use of the following shorthand expressions $\hat{\theta} := p(m)\beta$.

Steps for the if part of First Best (Proposition 5)

Lemma 5. *Suppose $1 - \check{\theta} \leq \theta \leq \hat{\theta}$. Then there is an optimal contract $(s, y_F, y_E) = (m, \theta, \theta)$.*

Lemma 6. *Suppose $\theta \geq \max\{1 - \bar{\theta}, \underline{\theta}\}$. Then, the optimal contract implies $y_k = \theta$.*

Lemma 7. *Suppose $\underline{\theta} > \theta > \hat{\theta}$. Then the optimal contract implies $y_k = \theta$.*

For $\theta < \underline{\theta}$, Full embracing is optimal (Part of Proposition 3)

Lemma 8. *If $\theta \leq \underline{\theta}$, there is an incentive-compatible contract with $y_F = \theta$ and $y_E \geq \theta$. If, in addition, $\theta < \hat{\theta}$ and first best is not feasible, agent F 's (DEC) has slack for $\theta < \underline{\theta}$. In that contract $s_0^* = -m$ until the first time E leads. Thereafter, $s^* = m$.*

Lemma 9. *The contract constructed in the proof of Lemma 8 is optimal in its region.*

²⁵This decomposition is more general and we use it beyond the $\theta = 0$ case. In Online Appendix G, we provide a general version of this decomposition.

²⁶The case in which $s = m$ and E 's (DEC) binds is already considered in the proof of the first statement.

To prove this lemma, we first make use of the preliminary results to simplify the class of contracts we need to consider. Then, we show that any contract in that class that does not coincide with the candidate from Lemma 8 cannot be optimal.

For $\theta \in (\underline{\theta}, \bar{\theta})$ Partial Embracing is Optimal (Part of Proposition 3)

Lemma 10. *If $\theta \in (\underline{\theta}, \bar{\theta})$, there exists a contract in which (DEC) binds for both agents and $y_F = \theta$. In that contract, $s = -m$ until the first time E leads. Thereafter, $s_E = m$ when E led last and $s_F < m$ when F led last.*

Lemma 11. *The contract constructed in the proof of Lemma 10 is optimal in its region if $\theta < \max\{1 - \bar{\theta}, 1 - \check{\theta}\}$.*

To prove optimality here, we exploit that players' (DEC) bind, leading to a well-defined set of Bellman equations.

For $\theta > \bar{\theta}$ Opportunism is Optimal (Proposition 4)

Lemma 12. *If $\theta > \bar{\theta}$, an optimal contract is such that $y_i \neq \theta$ and $s_E = -s_F = m$.*

To prove this lemma we combine the optimality result from the previous lemma with our usual construction approach.

Only-if part of Proposition 5

We have shown that for any region, other than those in Lemmas 5 to 7, the optimal contract does not achieve first best, which proves the only-if part.

C.2 Proof of Proposition 6

Proof. Proposition 6 follows from observing that y_E^* (y_F^*) decreases (increases) in θ for $\theta < \bar{\theta}$ and remains constant thereafter. $\square$

C.3 Proof of Proposition 7

Proof. We prove the following lemma, which implies Proposition 7.

Lemma 13. *Fix any (m, β) such that $2\beta m < 1 - \beta$ and assume $\theta = 1/2$. There exists a threshold $b^{FB} > 0$ such that for $b < b^{FB}$, the optimal commitment contract is the opportunistic contract, while for $b \geq b^{FB}$ the optimal commitment contract is the first-best contract. If, instead, $2\beta m \geq 1 - \beta$, first best is implementable for any $b > 0$ if $\theta = 1/2$.*

Proof. We prove the results in steps. First, we prove that the optimal contract at $\theta = 1/2$ is either the opportunistic contract or the first-best contract. Then we show that for $2\beta m < 1 - \beta$, a cutoff $b^{FB} > 0$ exists separating the two. We conclude by showing that for $2\beta m \geq 1 - \beta$ the first best is possible for any $b > 0$.

Step 1. $\bar{\theta}$ and $\check{\theta}$ at $1/2$. Recall that the first-best is implementable if and only if $\theta \geq \max\{1 - \bar{\theta}, 1 - \check{\theta}\}$. We will now show that at $\theta = 1/2$, for any (m, b, β) , the relevant first-best implementation constraint is $\theta \geq 1 - \bar{\theta}$. We do this by showing that if $1/2 \geq 1 - \bar{\theta}$,

then $1/2 > 1 - \check{\theta}$. To do that, notice that both $\check{\theta}$ and $\bar{\theta}$ are linearly increasing in b . Now fix an arbitrary $(m, \check{\beta}, b)$ such that $\check{\theta} = 1/2$; then, implicitly, it must hold that

$$\check{\beta} = \frac{1}{1 + 2m(2b + 1)} \Rightarrow \frac{\check{\beta}}{1 - \check{\beta}} = \frac{1}{2m(1 + 2b)}.$$

Replacing $\check{\beta}$ inside $\bar{\theta}$ gives us

$$\bar{\theta}_{\beta=\check{\beta}} = \frac{1 + b}{2(1 + 2b)} \frac{(1 + 2(1 + 4b)m)}{(1 + 2(1 + 2b)m)}$$

which is (weakly) larger than $1/2$ if and only if

$$(1 + b)(1 + 2(1 + 4b)m) \geq (1 + 2b)(1 + 2(1 + 2b)m) \Leftrightarrow m \geq 1/2.$$

which is impossible.

Step 2. Only Opportunistic or First-Best Optimal. Fixing $\theta = 1/2$ and selecting an arbitrary (m, β) , linearity in b implies that there exists a threshold $\check{b}$ such that $\theta < 1 - \check{\theta}$ for $b < \check{b}$ and $\theta > 1 - \check{\theta}$ for $b > \check{b}$. Notice further that at $b = \check{b}$, by construction, we have

$$\theta = 1 - \check{\theta} < 1 - \bar{\theta}.$$

The inequality comes from the order for $\check{\theta} = 1/2$ that we derived above. Equivalently, if $\theta = 1 - \bar{\theta} = 1/2$, then $b > \check{b}$. Finally, at $\theta = 1/2$ we must have either $\theta > \bar{\theta}$ or $\theta \geq 1 - \bar{\theta}$. Now, recall from Proposition 4 that if $\theta > \bar{\theta}$, the optimal contract is the opportunistic contract. As we see in the proof of Lemma 12, nothing in the argument relies on the fact that $\theta < 1/2$ and therefore the result extends to the limit. Thus, at $\theta = 1/2$, the optimal commitment contract is the opportunistic contract if $\bar{\theta} < 1/2$ and the first-best otherwise.

Step 3. Cutoff b^{FB} . Recall from Step 1 that $\bar{\theta}$ is strictly (and linearly) increasing in b and thus crosses $1/2$ from below once in b . Therefore, ignoring the non-negativity constraint for b , for any (m, β) , there exists some threshold $b' \in \mathbb{R}$ such that $\bar{\theta}$ is above $1/2$ if and only if $b > b'$. Moreover, if we fix (m, β) such that $2m\beta \geq 1 - \beta$, then $\lim_{b \rightarrow 0} \bar{\theta} \geq 1/2$, which means $b' \leq 0$ and thus $\bar{\theta}$ is strictly larger than $1/2$ for any $b > 0$ and first best is implementable for all positive b . If, instead, we fix (m, β) such that $2m\beta < 1 - \beta$, then $\lim_{b \rightarrow 0} \bar{\theta} < 1/2$, which implies the cutoff from the lemma, $b^{FB} > 0$, exists. $\square$

What remains is to show the first-period behavior in the opportunistic contract, which follows directly from the observation that $|\bar{\theta} - 1/2| = |(1 - \bar{\theta}) - 1/2|$ and thus the principal is indifferent between endorsing either agent in the first period. $\square$

D Penal Codes

Here, we provide further details on the optimal penal codes for the players. The results we state here provide the backbone of the punishment discussion in Section 4.2. Due to space restrictions, we defer their proofs to Online Appendix G.

Note first that *simple penal codes* (as in Abreu (1988), Proposition 5) suffice in our game due to perfect information. The culprit is always observed by other players, and so is their “crime,” which makes more complicated punishment schemes, as in, e.g., Mailath et al. (2017), not relevant. Existence is guaranteed by the compactness of the equilibrium payoff set.

Note: By self-generation of equilibrium payoff sets, penal codes are nothing other than those contracts that minimize the punished player’s ex-ante value. We will therefore adjust the notation here to reflect this. Specifically, we consider h_0 to be the history at which the punishment started, i.e., the node at which the deviation occurred.²⁷

D.1 The Principal’s Penal Code

We show that a stationary penal code for the principal exists.

Lemma 14. *There is a penal code for the principal in which $y_E = 1$ whenever E leads.*

The intuition is straightforward: If there was a penal code in which $y_E < 1$, then such a contract can only be an optimal penal code for the principal if the principal promises E some endorsement in return. But then, using the logic of Lemma 2 in reverse, we can find a contract in which E selects $y_E = 1$. Having $y_E = 1$ established, we can now construct the optimal penal code for the principal.

Lemma 15. *A stationary optimal penal code for P exists. It is characterized by:*

- *the principal endorses F , $s = -m$,*
- *E fully polarizes $y_E = 1$*
- *F either fully polarizes, $y_F = 0$, or chooses $y_F = \hat{y}_F$, where $\hat{y}_F$ is such that his dynamic enforcement constraint binds.*

Proof. From Lemma 14, we have $y_E = 1$. Hence, irrespective of $y_F(\cdot)$, P is worse off if E leads because $\theta < 1/2$. Thus, P chooses $s = -m$. To conclude, we characterize F ’s actions. Fixing $(s = -m, y_E = 1)$, let $\hat{y}_F$ be F ’s policy such that his (DEC) binds. Note further that if a certain $\hat{y}_F$ satisfies (DEC) for some s , it satisfies (DEC) for $s = -m$. It is straightforward to see that F ’s action is either $y_F = 0$, which is trivially incentive compatible, or $y_F = \hat{y}_F$, as P prefers any intermediate $y_F \in (0, \hat{y}_F)$ to at least one of these two extremes. If $\hat{y}_F < 2\theta$, $y_F = 0$ is worse for P , and otherwise, $y_F = \hat{y}_F$ is worse. $\square$

D.2 The Friend’s Penal Code

Note. Compactness of the equilibrium payoff set and self-generation imply that for every implementable continuation value for P , w_p , there is a contract that minimizes F ’s continuation value conditional on yielding w_P to P . The key to the construction of the Friend’s penal code is the back-to-business property, which we now define formally.

Definition 3 (Back-to-business property). A penal code for the Friend, C^F , has the *back-to-business property*, if any on-path history in that contract, $h \in \mathcal{H}_F^*(C^F)$ implies a

²⁷We sometimes use $\underline{w}_k(h_0)$ instead of $\underline{w}_k$ to emphasize that this is the *ex-ante* payoff to k .

continuation game identical to that of a history $h \notin \mathcal{H}_F(C^*)$ in the optimal contract C^* that selects F as the leader.

We now show that a penal code for the Friend that sustains the optimal contract has the back-to-business property. The first case we address is when F 's on-path (DEC) binds. Since F 's (DEC) binds, a penal code that implements the optimal contract has to be optimal. We see that this optimal penal code has the back-to-business property.

Lemma 16. *Suppose F 's (DEC) binds in the embracing phase of the optimal contract and first best is infeasible. The optimal penal code for F has the back-to-business property.*

We prove Lemma 16 in Online Appendix G. The construction admits Corollary 1.

Corollary 1. *The optimal penal code for F consists of two phases: (i) A punishment phase that ends when F leads for the first time since the beginning of the penal code, and (ii) the back-to-business phase.*

1. *If the optimal contract is such that P has slack in the embracing phase, then the punishment phase is such that the principal chooses a full embracing strategy $s_0^F = m$, and the Enemy chooses $y_E^F = \min\{\hat{y}_E, 1\}$, where $\hat{y}_E$ is such that the principal's (DEC') holds with equality.*
2. *If the optimal contract is such that P 's (DEC') holds with equality, the punishment phase replicates the on-path behavior for histories in which E was the last to lead.*

The idea is straightforward: because F 's (DEC) binds in the embracing phase, this contract is already the worst for F conditional on leading. In the initial punishment phase, the result relies on the idea familiar from the proof of Proposition 2: If all policy choices $y_k \geq \theta$, we can decompose the Friend's payoff into a policy payoff and a leadership rent. Thus, by ensuring that P 's (DEC') binds, we ensure that indeed we are minimizing F 's ex-ante payoffs. If P 's (DEC') has slack even for $y_E = 1$, we have recreated the commitment punishment (in value) for F . Our next result complements the above, showing that a penal code with back-to-business is always a valid penal code for implementing the full-embracing optimal contract.

Lemma 17. *Whenever the optimal contract features full embracing in the embracing phase, it can be supported by a penal code for F , C^F , which consists of (i) a punishment phase that ends when F leads for the first time since the beginning of the penal code, and (ii) the back-to-business phase. Moreover, in the punishment phase:*

- $s = m$
- $y_E = \min\{\hat{y}_E, 1\}$, where $\hat{y}_E \in [0, \infty)$ is chosen such that P 's (DEC') binds.

The result relies on the decomposition (since $y_k \geq \theta$). The leadership rent is minimized under $s^F = m$. The policy payoff is minimized by making P 's (DEC') bind at the initial node. If P has slack even for $y_E = 1$, either F 's (DEC) binds, in which case we have the commitment punishment, or it has slack, in which case the penal code is not necessarily optimal, but still valid to sustain the optimal contract.

Conclusion. We demonstrate here that the penal codes for P and F are linked through F 's action when punishing P , and the back-to-business property of F 's penal code. Thus, they jointly form a system of equations that behaviorally characterizes both punishments and the optimal contract. However, depending on details, different constraints hold with equality, making a simple closed-form characterization difficult. In the supplementary material, we provide code that obtains the full symbolic characterization of all penal codes and the optimal contract for any arbitrary parameter constellation.

E The Optimal No-Commitment Contract

E.1 Proof of Proposition 8

Proof. For existence of $\bar{b}(\theta)$, observe that as $b \rightarrow \infty$, both $\check{\theta}, \bar{\theta} \rightarrow \infty$; thus, with commitment, P could implement the first best, in particular via the following strategy: fully endorsing the agent who has led last.

We will now show that the first-best can be implemented without commitment when b is large. For that purpose, assume the following punishment: If an agent deviates, the principal fully endorses the non-deviator for one period before returning to the on-path contract. Agents continue to choose $y_k = \theta$. That punishment is sufficient to sustain $y_k = \theta$ if $2\beta m(b - |\theta_k - \theta|) \geq |\theta_k - \theta|$ which holds for $b \geq (|\theta_k - \theta|)(1 + 2\beta m)/(2\beta m)$. P 's (DEC') holds on-path and off-path because all agents are expected to always select θ .

The existence of $\hat{\theta}$ follows from Lemma 18 which we state and prove below. $\square$

Lemma 18. *Fix (m, b, β) . There exists a threshold $\hat{\theta} < 1/2$ such that for any $\theta \in [\hat{\theta}, 1/2)$, even if the principal has no commitment power, she optimally implements a contract that is observationally equivalent to the optimal commitment contract.*

Proof. Note that, by construction, the principal prefers the on-path contract to the NES payoff. Moreover, the Enemy's grim-trigger punishment is the NES contract, which is (trivially) implementable also absent principal commitment. Thus, if the principal can sustain a punishment that gives the Friend the same off-path payoff as the grim-trigger punishment used in the construction of the commitment contract, the principal can implement a contract observationally equivalent to the commitment contract. We now construct such punishments.

At the Limit. At the limit $\theta = 1/2$, either first best or the opportunistic contract is the commitment optimum (by Lemma 13). Now, observe that the following class of contracts is also implementable at the limit: Parties choose their bliss points, $y_F = 0$ and $y_E = 1$, whenever they lead, and the principal chooses an arbitrary strategy $s \in [-m, m]$ at *any* history. To see that, consider the static game: incentives for F and E hold trivially, and both policy choices lead to a payoff of $-1/2$ for the principal. That makes P indifferent between *any* s . Because the repetition of a stage Nash equilibrium is also an equilibrium

of the repeated game, the claim follows. Note that this claim means in particular that the grim-trigger punishment for F under commitment is implementable at the limit $\theta = 1/2$.

Outside the limit. We now move slightly away from the limit $\theta = 1/2$ and consider $\theta < 1/2$ but close to it. We want to implement the optimal on-path commitment contract $(s_F^*, s_E^*, y_E^*, y_F^*)$ and consider the following punishment contract for the Friend: If F deviates, P supports E , who does not moderate at all ($y_E^F = 1$) until the Friend returns to the lead. From that point onward the continuation game follows the on-path commitment contract. If the principal or the Enemy deviates, the continuation game is NES.

The principal's value from punishing the Friend is

$$w_P^F = (1/2 + m)(-(1 - \theta) + \beta w_P^F) + (1/2 - m)(v_P^*(F)) = -\frac{p(m)(1 - \theta) - (1 - p(m))v_P^*(F)}{1 - p(m)\beta}$$

If the principal deviates her payoff is

$$\underline{w}_P = w_P^E = -\frac{p(m)\theta + (1 - p(m))(1 - \theta)}{1 - \beta}.$$

Now at the limit $\underline{w}_P(\theta = 1/2) = -1/(2(1 - \beta))$, and $w_P^* > \underline{w}_P$ because if $\bar{\theta} \leq 1/2$ first best is implementable, and otherwise $y_F = \bar{\theta}$, $y_R = 1 - \bar{\theta}$ is implementable (Lemma 13). Both are better than NES. Moreover, notice that

$$v_P^*\left(F; \theta = \frac{1}{2}\right) \geq -\frac{1}{2} + \beta w_P^*\left(F; \theta = \frac{1}{2}\right) > -\frac{1}{2} + \beta \underline{w}_P\left(\theta = \frac{1}{2}\right) = -\frac{1}{2(1 - \beta)} = \underline{w}_P\left(\theta = \frac{1}{2}\right).$$

But then, we have that $w_P^F(\theta = 1/2) > \underline{w}_P(\theta = 1/2)$ which implies strict slack of the principal's incentives at the limit $\theta = 1/2$. Since both $w_P^F(\theta)$ and $\underline{w}_P(\theta)$ are continuous in θ , the principal's incentives are maintained even outside the limit.

That implies that the commitment contract is sustainable for some θ outside the limit. Finally, if that punishment is sustainable for the principal, it is sustainable for both parties. The Enemy gains from choosing 1, and the Friend is punished by being threatened to go through a punishment period of $y_E^P = 1$ and $s_F^P = m$ instead of $s_F^* \leq m$ and $y_E^* < 1$. Thus, both parties' (DEC) and the principal's (DEC') hold and the no-commitment contract is observationally equivalent to the commitment contract for θ close to $1/2$. $\square$

E.2 Proof of Proposition 9

Proof. To prove Proposition 9, we have to show that whenever a region exists with commitment, part of that region survives removing commitment. We do this by first showing that as $\theta \rightarrow 1/2$, full embracing vanishes for any b if there ever is a partial embracing or opportunistic region. Moreover, as $\theta \rightarrow 1/2$, partial embracing survives for intermediate b . We then use the ‘‘centrism is commitment’’ result of Proposition 8 to show robustness of parts of the first best, the partial embracing, and the opportunistic region.

Then, we consider the other corner $\theta \rightarrow 0$ and show that full embracing survives in some neighborhood of $\theta = 0$ if it survives at $\theta = 0$.

Lemma 19. *If $1/2 > \theta > \theta^\dagger := \beta(1/2 + m)$, there exist thresholds $b^\dagger < b^\ddagger$ such that the optimal commitment contract is the opportunistic contract for $b < b^\dagger$, the first best for $b > b^\ddagger$, and a partial embracing contract for $b \in (b^\dagger, b^\ddagger)$.*

Proof. We prove this statement by fixing (m, β) and showing two steps separately.

1. showing that for $\theta = \theta^\dagger$, there is a unique $b^\diamond$ such that $1 - \bar{\theta} = 1 - \check{\theta} = \underline{\theta} = \theta^\dagger$,
2. Fixing (m, β) and $\theta \in (\theta^\dagger, 1/2)$:
 - (a) if b is such that $\theta \leq \underline{\theta}$, then $\theta > \max\{1 - \bar{\theta}, 1 - \check{\theta}\}$.
 - (b) if b is such that $\theta \geq \max\{1 - \bar{\theta}, 1 - \check{\theta}\}$ then $\theta > \bar{\theta}$
 - (c) if b is such that $\theta \geq \bar{\theta}$, then $\theta < \max\{1 - \bar{\theta}, 1 - \check{\theta}, \underline{\theta}\}$.

Result (1) implies that the three cutoffs coincide; then result (2) implies that above their intersection (in the θ space), a full embracing contract does not exist (2(a)), but a partial embracing contract (2(b)) and an opportunistic contract (2(c)) exist for some b .

Result (1). First, recall that all three cutoffs, $\underline{\theta}, \bar{\theta}, \check{\theta}$, are linearly increasing in b . Thus any pair of $1 - \bar{\theta}, 1 - \check{\theta}, \underline{\theta}$ intersects at most once in b . The following Lemma states that they intersect at the same point. We prove it in Online Appendix G.

Lemma 20. *Assume $2m\beta < 1 - \beta$. The following are equivalent for $\theta^\dagger = \beta(1/2 + m)$: (1) $b = b^\diamond := \frac{1-\beta}{2\beta m} - 1$, (2) $\underline{\theta} = \theta^\dagger$, (3) $1 - \bar{\theta} = \theta^\dagger$, (4) $1 - \check{\theta} = \theta^\dagger$.*

Result (2). Note that

$$0 < \frac{\partial \check{\theta}}{\partial b} = 2m\beta < 2m\beta \underbrace{\frac{1 - \beta(1/2 + m)}{1 - \beta}}_{>1} = \frac{\partial \bar{\theta}}{\partial b}$$

which implies both $1 - \check{\theta}$ and $1 - \bar{\theta}$ are decreasing in b with $1 - \bar{\theta}$ decreasing slower. That implies if $\theta = 1 - \bar{\theta} > \theta^\dagger$ then $\theta > 1 - \check{\theta}$. Because $\underline{\theta}$ increases in b and intersects with $1 - \bar{\theta}$ at $(b^\diamond, \theta^\dagger)$, the 2(a) follows. The two other results, 2(b) and 2(c), follow because by construction $\bar{\theta} = 1 - \bar{\theta}$ at $(b^{FB}, 1/2)$ (see Proposition 7), and for $1/2 > \theta = \bar{\theta} < 1 - \bar{\theta}$. The last result, 2(c), follows because $\bar{\theta}$ increases in b and intersects with $1 - \bar{\theta}$ at $(b^{FB}, 1/2)$, and for $1/2 > \theta = \bar{\theta} < 1 - \bar{\theta}$. $\square$

Lemma 19 combined with Proposition 8 proves robustness of a partial embracing region, a first-best region, and an opportunistic region if they exist in the first place.²⁸

Lemma 21. *Fix (m, β) and pick $b \in (\bar{b}_0, \check{b})$.²⁹ Then there exists $\theta^\P > 0$ such that for any $\theta \in (0, \theta^\P)$, the optimal no-commitment contract features full embracing and is observationally equivalent to the optimal commitment contract.*

Proof. Because $\underline{\theta} > 0$, for any (m, β, b) , there exists a region where full embracing is optimal under commitment. Also, if $b < \check{b}$, then there exists a $\theta' > 0$ such that for $\theta \in [0, \theta')$,

²⁸ $\theta^\dagger > 1/2 \Leftrightarrow 2m\beta > 1 - \beta$ which implies that only full embracing and first best are commitment optima.

²⁹These cutoffs are defined in Proposition 2 and Lemma 13 respectively.

the optimal commitment contract is not first best.³⁰ Consider now the following results for $\theta = 0$. First, for any $b \in (\bar{b}_0, \check{b})$, the optimal commitment contract is implementable without commitment by Proposition 2 and because $b > \bar{b}_0$, P 's (DEC') has slack and $\hat{y}_E > y_E^*$, where $\hat{y}_E$ is E 's action during the punishment phase of F 's penal code (see the proof of Proposition 2 for details). Second, because $\hat{y}_E > y_E^*$ and F 's punishment has the back-to-business property (see Lemma 17), F 's (DEC) also has slack.

Now, with these observations, the result follows from continuity in θ of the commitment continuation payoffs $w_P^*(h)$, $w_F^*(h)$, and the penal codes for the non-commitment case, $\underline{w}_F$, and $\underline{w}_P$. Continuity in the on-path payoffs follows directly from the characterization of the optimal commitment contract (Appendix C), while continuity in P 's punishment payoffs follows from the construction in Lemmas 14 and 15. Finally, continuity in F 's punishment payoff follows from the back-to-business property, the continuity of the on-path contract, and continuity of $\hat{y}_E$ in θ , which is a consequence of the continuity of $\underline{w}_P$. $\square$

Lemma 21 shows that if $b > \bar{b}_0$, then there is a region (with $\theta > 0$) in which commitment plays no role and full embracing survives in the absence of commitment. $\square$

F Implications

F.1 Proof of Proposition 10

For the commitment case, the result follows from the observation that $\underline{\theta}$ is monotone in all three parameters.

For the no-commitment case, take F 's (DEC) binding in the full embracing contract:

$$y_F^* = \theta = \beta(w_F^* - \underline{w}_F) \quad (6)$$

Recall that F 's optimal penal code involves back-to-business (by Lemmas 17-16), and thus both the optimal contract and F 's punishment have all $y_k \geq \theta$, which implies one can express F 's payoffs as P 's payoffs plus his expected leadership rents:

$$w_F^* = \frac{1-p(m)}{1-\beta}b + w_P^* - \frac{\theta}{1-\beta} \quad \text{and} \quad \underline{w}_F = \frac{1-p(m)}{1-\beta}b + w_P^F - \frac{\theta}{1-\beta},$$

where $w_P^F = \underline{w}_P$ in the relevant region, since otherwise F 's commitment punishment is attainable. Substituting in (6), we get $\theta = \beta(w_P^* - \underline{w}_P)$. As $\theta > 0$, P 's (DEC') has slack on path, and the frontier of the full embracing region is given by (6). The result follows from the fact that, for either penal code of P , $\beta(w_P^* - \underline{w}_P)$ increases both in β and b .

F.2 Proof of Proposition 11

Proof. For the commitment case, Proposition 11 follows for $\theta < \bar{\theta}$ because $y_F^* = \theta$ and y_E^* (weakly) decreases in θ . For $\theta \geq \bar{\theta}$, both y_F^* and y_E^* are constant, implying that the principal loses in the exclusion phase as θ increases; in the embracing phase, the changes cancel each other out.

³⁰Details are found in the proof of Lemma 19.

Lemma 18 implies that for θ close to $1/2$ commitment plays no role, thus the principal's ex-ante payoff decreases in θ if $\bar{\theta} < 1/2$ also without commitment. This is straightforward for $b < \bar{b}_0$, in which case the no-commitment payoff is that of NES (by Proposition 2), which is dominated by the opportunistic contract under commitment for θ close to $1/2$.

For the remaining case, the following lemma (which we prove in Online Appendix G) implies the result in the no-commitment case for $b > \bar{b}_0$.

Lemma 22. *There exists a $b_\Delta < \bar{b}_0$ such that for $b > b_\Delta$, the principal's payoff under commitment at $\theta = 0$ is smaller than at $\theta \rightarrow 1/2$.*

Since the statement holds for $b \geq \bar{b}_0 > b_\Delta$, this completes our proof. $\square$

F.3 Proof of Proposition 12

Proof. The first part of Proposition 12 follows because when firing the Enemy, the principal receives a per-period payoff of $-\theta$. In the optimal contract, the principal's payoff is continuous and she strictly prefers it over firing at the boundaries $\theta \in \{1/2, 1 - \bar{\theta}, 1 - \check{\theta}\}$, one of which is always admissible. Because $\Upsilon(b, m)$ is non-empty and payoffs are continuous, the result follows. For the second part of Proposition 12, observe that for a given θ the principal weakly prefers competition if $y_E^* - \theta \leq \theta$. Because commitment does not matter close to $\theta = 1/2$, the result continues to hold.

Now, if, for the chosen (b, m, θ) , the first best is achieved for some $\hat{\beta}$ then, at $\hat{\beta}$, the left-hand side of the inequality is 0. Continuity and monotonicity of y_E^* in β imply that there is a $\tilde{\beta} < \hat{\beta}$ such that competition is preferred. Observe that $\bar{\theta} \rightarrow \infty$ as $\beta \rightarrow 1$ and $\check{\theta} \rightarrow 1/2 + m(1 + 2b) \geq 1$ if and only if $2b \geq \frac{1-2m}{2m}$. Thus, under the condition, first best is achieved at least at the limit, which implies that, for every θ , there is a $(\hat{\beta}, \theta) \in \Upsilon(b, m)$ such that the principal does not want to fire the agent. Because the first best survives no commitment, the result follows in the no-commitment case. $\square$

References

- Abreu, D. (1986). “Extremal equilibria of oligopolistic supergames”. *JET*, pp. 191–225.
- (1988). “On the theory of infinitely repeated games with discounting”. *Econometrica*, pp. 383–396.
- Acemoglu, D., M. Golosov, and A. Tsyvinski (2011). “Power fluctuations and political economy”. *JET*, pp. 1009–1041.
- Acharya, A., E. Lipnowski, and J. Ramos (2024). “Political Accountability Under Moral Hazard”. *AJPS*, forthcoming.
- Albarelo, A. (2024). “Coalition policy in multiparty governments: whose preferences prevail”. *Political Science Research and Methods*, pp. 318–335.

- Alonso, R., W. Dessein, and N. Matouschek (2008). “When Does Coordination Require Centralization?” *AER*.
- Alonso, R. and N. Matouschek (2008). “Optimal delegation”. *ReStud*, pp. 259–293.
- Andrews, I. and D. Barron (2016). “The allocation of future business: Dynamic relational contracts with multiple agents”. *AER*, pp. 2742–2759.
- Anesi, V. and P. Buisseret (2022). “Making elections work: Accountability with selection and control”. *AEJ:Micro*, pp. 616–44.
- Barron, D. and M. Powell (2019). “Policies in relational contracts”. *AEJ:Micro*, pp. 228–49.
- Bendor, J., A. Glazer, and T. Hammond (2001). “Theories of delegation”. *Annual review of political science*, pp. 235–269.
- Bendor, J. and A. Meirowitz (2004). “Spatial models of delegation”. *APSR*, pp. 293–310.
- Bichay, N. (2026). “Converging to the mean? Moderating ideologies of far-right and far-left parties in government”. *Party Politics*, p. 13540688251339632.
- Board, S. (2011). “Relational contracts and the value of loyalty”. *AER*, pp. 3349–3367.
- Bolton, P. and M. Dewatripont (2013). “Authority in organizations”. *Handbook of organizational Economics*, pp. 342–372.
- Callander, S. (2008). “A theory of policy expertise”. *QJPS*, pp. 123–140.
- Calzolari, G. and G. Spagnolo (2020). “Relational contracts, procurement competition, and supplier collusion”. *mimeo*.
- Deb, J., A. Kuvalekar, and E. Lipnowski (2025). “Fostering Collaboration”. *Theoretical Economics*, forthcoming.
- Delgado-Vega, Á. (2025). “Which Side are You on? Interest Groups and Relational Contracts”. *mimeo*.
- Dessein, W. (2002). “Authority and Communication in Organizations”. *ReStud*.
- Dixit, A., G. M. Grossman, and F. Gul (2000). “The dynamics of political compromise”. *JPE*, pp. 531–568.
- Gibbons, R. (2020). “March-ing toward organizational economics”. *Industrial and Corporate Change*, pp. 89–94.
- Gibbons, R. and K. J. Murphy (1992). “Optimal Incentive Contracts in the Presence of Career Concerns: Theory and Evidence”. *JPE*, pp. 468–505.
- Ginsborg, P. (1990). *A History of Contemporary Italy: 1943-80*. Penguin UK.
- Invernizzi, G. M. and M. M. Ting (2024). “Institutions and political restraint”. *AJPS*, pp. 58–71.
- Leonardi, R. and D. A. Wertman (1989). *Italian Christian Democracy: the politics of dominance*. Springer.
- Levin, J. (2002). “Multilateral contracting and the employment relationship”. *QJE*, pp. 1075–1103.
- (2003). “Relational incentive contracts”. *AER*, pp. 835–857.
- Li, J., N. Matouschek, and M. Powell (2017). “Power dynamics in organizations”. *AEJ:Micro*, pp. 217–241.
- Lipnowski, E. and J. Ramos (2020). “Repeated delegation”. *JET*, p. 105040.
- Luo, Z., Z. Liu, Y. Qiu, and S. Yu (2025). “A Relational Theory of Power Alternation”. *Available at SSRN*.
- MacLeod, W. B. (2014). *The Economics of Relational Contracts*.

- Maier, C. S. (2015). *Recasting bourgeois Europe: stabilization in France, Germany, and Italy in the decade after World War I*. Princeton University Press.
- Mailath, G. J., V. Nocke, and L. White (2017). “When and How the Punishment Must Fit the Crime”. *International Economic Review*.
- Malcomson, J. M. (2012). “Relational Incentive Contracts”. *The Handbook of Organizational Economics*. Princeton University Press, pp. 1014–1065.
- March, J. G. (1962). “The Business Firm as a Political Coalition”. *The Journal of Politics*, pp. 662–678.
- Martin, L. W. and G. S. Vanberg (2004). “Policing the Bargain: Coalition Government and Parliamentary Scrutiny”. *AJPS*, pp. 13–27.
- McClellan, A. and X. Liao (2023). “Intraparty Politics and Dynamic Policy Polarization”. *mimeo*.
- Nocke, V. and R. Strausz (2023). “Collective Brand Reputation”. *JPE*, pp. 1–58.
- Peukert, D. (1993). *The Weimar Republic*. Macmillan.
- Rayo, L. (2007). “Relational incentives and moral hazard in teams”. *ReStud*, pp. 937–963.
- Tirole, J. (1994). “The Internal Organization of Government”. *Oxford Economic Papers*. New Series, pp. 1–29.
- Van Spanje, J. and W. Van Der Brug (2007). “The party as pariah: The exclusion of anti-immigration parties and its effect on their ideological positions”. *West European Politics*, pp. 1022–1040.
- Watson, J. (2021). “Theoretical foundations of relational incentive contracts”. *Annual Review of Economics*, pp. 631–659.
- Wright, J. (2002). *Gustav Stresemann: Weimar’s greatest statesman*. Oxford University Press, USA.

Data Availability Statement

The code used to generate Figures 2 and 3 in this paper is available at <https://doi.org/10.5281/zenodo.21359697>.