

The Political Economy of Indirect Control*

Gerard Padró i Miquel[†] Pierre Yared[‡]

October 31, 2009

Abstract

This paper characterizes optimal policy in an environment in which a government uses indirect control to exert its authority. We develop a dynamic principal-agent model in which a principal (a government) delegates the prevention of a disturbance—such as riots, protests, terrorism, crime, or tax evasion—to an agent who has an advantage in accomplishing this task. Our setting is a standard dynamic principal-agent model with two additional features. First, the principal is allowed to exert direct control by intervening with an endogenously determined intensity of force. Second, the principal suffers from limited commitment. Using recursive methods, we derive a fully analytical characterization of the likelihood, intensity, and duration of intervention in the optimal contract. The first main insight from our model is that repeated and costly interventions are a feature of optimal policy. This is because they serve as a punishment to induce the agent into desired behavior. The second main insight is a detailed analysis of a fundamental tradeoff between the intensity and duration of intervention which is driven by the principal’s inability to commit. Finally, we derive sharp predictions regarding the impact of various factors on the optimal likelihood, intensity, and duration of intervention.

Keywords: Institutions, Asymmetric and Private Information, Structure of Government, Dynamic Contracts

JEL Classification: D02, D82, H1

*We would like to thank Daron Acemoglu, Mike Golosov, Johannes Horner, Narayana Kocherlakota, Nancy Qian, Kjetil Storesletten, Aleh Tsyvinski, and seminar participants at Kellogg MEDS for comments.

[†]London School of Economics and NBER. email: g.padro@lse.ac.uk.

[‡]Columbia University. email: pyared@columbia.edu.

1 Introduction

In exerting their authority, governments often use indirect control: Certain political responsibilities are left to local agents or warlords who have an advantage in fulfilling them. These tasks range from the provision of law and order, the prevention of riots and protests, the control of terrorism and insurgency, to the collection of taxes. For example, by the first century, the Romans had established a series of client states and chieftaincies along their borders which gave them control of a vast territory with great economy of force. These clients were kept in line by a combination of subsidies and favors and threatened by occasional military intervention.¹ Beyond Roman times, this strategy of indirect control through local agents has been used by the British during colonial times and the Turks during the Ottoman era, and it is tacitly used today by many governments.² This suggests the following question: How should a government use rewards and interventions to align the incentives of the local agent with its own?

In answering this question, it is important to take into account that the interaction between a government and a local agent is inherently dynamic, and that there are three key political economy frictions to consider. First, the government cannot commit to providing rewards or using interventions. Second, the local agent cannot commit to fulfilling his delegated task. Third, the local agent's actions, which often occur through informal channels, are imperfectly observed by the government. The optimal policy in this context must take into account the interaction between double-sided lack of commitment and asymmetric information. As such, a modified dynamic principal-agent model (in which the government is a principal) can provide guidance on the implications of these frictions for optimal policy.

In this paper, we develop such a model. The principal delegates the prevention of a disturbance—such as riots, protests, terrorism, crime, or tax evasion—to an agent who has an advantage in accomplishing this task. Our setting is a standard dynamic principal-agent model with two additional features which are natural in our application.³ First, the

¹See Syme (1933) and Luttwak (1976).

²This is particularly the case in governments that have tenuous control over parts of their territory, for instance, in Pakistan's Federally Administered Tribal Areas and in rural areas in many African countries. On this point, see Herbst (2000) and Reno (1998). Recent interventions such as Pakistan in its tribal territories, Russia in Chechnya, Israel in the Palestinian Territories, or Indonesia in Banda Aceh arguably fit the pattern. The United Kingdom also suspended local administration and deployed the army during The Troubles in Northern Ireland from 1968 to 1998.

³The literature on dynamic principal-agent relationships is vast and cannot be summarized here. Some examples are Acemoglu, Golosov, and Tsyvinski (2008), Albuquerque and Hopenhayn (2002), Atkeson and Lucas (1992), Fong and Li (2009), Golosov, Kocherlakota, and Tsyvinski (2003), Phelan (1995), and Thomas and Worrall (1990). Also see Debs (2009), Egorov and Sonin (2009), and Myerson (2008) for

principal is allowed to exert direct control by intervening with an endogenously determined intensity of force. Second, the principal suffers from limited commitment. We focus on characterizing the optimal likelihood, intensity, and duration of intervention. Using the recursive methods of Abreu, Pearce, and Stacchetti (1990), we derive a fully analytical characterization of the optimal contract. The first main insight from our model is that repeated and costly interventions are a feature of optimal policy. This is because they serve as a punishment to induce the agent into desired behavior.⁴ A second insight, which emerges from our explicit characterization, is the existence of a fundamental tradeoff between the intensity and duration of intervention that is driven by the principal's inability to commit. Finally, we derive sharp predictions regarding the impact of various factors on the optimal likelihood, intensity, and duration of intervention.

More specifically, we construct a repeated game between a principal and an agent where in every period, the principal decides whether or not to intervene. Under intervention, he chooses the intensity of force, where higher intensity is costly to both the agent and the principal (i.e., it does not help to reduce the probability of a disturbance) and features diminishing returns (i.e., the marginal pain inflicted on the agent is decreasing in intensity). The principal cannot commit to future actions. If the principal does not intervene, the agent can reduce the probability of disturbances by exerting unobservable effort which can be high or low. Both players are strictly better off under high effort by the agent compared to intervention by the principal. Nonetheless, there are two limitations to the extent to which intervention can be avoided. First, the agent cannot commit to high effort once the threat of intervention has subsided. Second, the principal does not observe the agent's effort, and since disturbances might happen even under high effort, the agent can always unobservably deviate and pretend to have exerted high effort. Therefore, the Nash equilibrium of the stage game is intervention with minimal force (i.e., direct control). We consider the efficient sequential equilibrium of this game in which reputation sustains equilibrium actions, and we fully characterize in closed form the long run dynamics of the

applications to delegation problems in dictatorships.

⁴The use of costly interventions as punishment is very common in situations of indirect control. In his discussion of the Ottoman Empire, Luttwak (2007) writes:

"The Turks were simply too few to hunt down hidden rebels, but they did not have to: they went to the village chiefs and town notables instead, to demand their surrender, or else. A massacre once in a while remained an effective warning for decades. So it was mostly by social pressure rather than brute force that the Ottomans preserved their rule: it was the leaders of each ethnic or religious group inclined to rebellion that did their best to keep things quiet, and if they failed, they were quite likely to tell the Turks where to find the rebels before more harm was done." (p.40)

optimal contract.

Our first result is that repeated and costly interventions are a feature of optimal policy. Specifically, the optimal contract after a sufficient number of disturbances features two phases of play: a cooperative phase and a punishment phase that sustain each other. In the cooperative phase, the agent exerts high effort because he knows that a disturbance can trigger a transition to the punishment phase. In the punishment phase, the principal temporarily intervenes with a unique endogenous level of intensive force. The principal exerts costly force because failure to do so triggers the agent to choose low effort in all future cooperative phases, making direct control—i.e., permanent intervention with minimal intensity—a necessity. Importantly, the optimal contract which maximizes the principal’s welfare under cooperation also minimizes the agent’s welfare under punishment. This is because conditional on the agent exerting high effort, the optimal policy must minimize the likelihood of punishment. To keep the agent’s incentive constraint satisfied, minimum likelihood is achieved by providing the worst feasible payoff to the agent in the punishment phase.

Our second result follows from our explicit characterization of the worst feasible punishment. Recall that the principal cannot commit to future actions. As a consequence, he can always deviate to permanent direct control, which constitutes his min-max payoff. This generates an incentive compatibility constraint on the side of the principal that produces a fundamental tradeoff between the duration and the intensity of credible interventions. In particular, he can only be induced to intervene with costly intensity if cooperation is expected to resume in the future, and higher intensity is only incentive compatible if cooperation resumes sooner. This link between intensity and duration generates a non-monotonic relationship between intensity and the agent’s welfare under punishment. At low levels of intensity, the agent’s welfare naturally declines when intensity rises. However, at higher levels of intensity, diminishing returns set in and the counteracting effect of shorter duration makes his expected welfare actually increasing in intensity. Since the principal seeks to minimize the agent’s welfare under punishment, it follows that there is a unique and interior level of intensity that is used.

Our final result concerns the effect of three important factors on the optimal likelihood, intensity, and duration of intervention. First, we consider the effect of a decline in the cost of intensity to the principal. Second, we consider the effect of a rise in the cost of disturbances to the principal. Finally, we consider the effect of a rise in the cost of effort to the agent.

We show that all three changes increase the optimal intensity and decrease the optimal duration of intervention. In the first case, it is clear that a reduction in the marginal cost

of intensity increases its optimal use. In the second case, as the cost of disturbances rises, so do the returns to leveraging the comparative advantage of the agent. As the prospect of direct control becomes worse, the principal is willing to raise the intensity of intervention. In the third case, as the cost of effort for the agent rises, higher levels of intensity become necessary to satisfy the agent's incentive constraint. In all three cases, due to the principal's incentive constraint, these increases in the level of intensity necessitate a decline in the duration of intervention.

Though all three changes increase the optimal intensity and decrease the optimal duration of intervention, only the third also raises its likelihood. Specifically, if the cost of intensity to the principal declines or if the cost of disturbances to the principal rises, then harsher punishments are feasible. Because the agent's incentive compatibility constraint is slackened by these changes, such punishments can be applied less often without weakening incentives for the agent. Therefore, the likelihood of intervention declines. In contrast, if the cost of effort to the agent rises, then incentives are harder to provide for the agent, and the likelihood of intervention must rise following the realization of a disturbance.⁵

As an aside, note that our benchmark model ignores three additional issues. First, it ignores the possibility that permanent concessions by the principal can reduce the presence of disturbances in the future. Second, it ignores the possibility that the agent's identity can change over time because of political transitions. Third, it ignores the possibility that high intensity levels by the principal today can raise the cost of effort by the agent in the future, for example if the agent becomes more antagonistic. These issues are discussed in our extensions which show that our main conclusions are unchanged.

This paper makes three contributions. First, it contributes to the dynamic principal-agent literature described in footnote 3 by allowing for costly intervention by a principal who suffers from limited commitment. Specifically, our model has the same structure as Fong and Li (2009) who also consider the effect of limited commitment by the principal in a labor market setting, though in contrast to their work we allow for costly intervention. Given that we study a model with double-sided lack of commitment, our results are related to the literature on repeated games with imperfect monitoring and to the insights due to the seminal work of Green and Porter (1984). These authors present examples of sequential equilibria with symmetric strategies in which punishment in the form of temporary price wars sustain cooperation between oligopolistic firms. In contrast to this work, we use the recursive methods of Abreu, Pearce, and Stacchetti (1990) to explicitly

⁵In other words, when the cost of effort increases, the principal uses two margins to adjust punishments. He increases both intensity and likelihood to meet the tighter incentive compatibility constraint of the agent.

characterize the efficient equilibrium under history-dependent strategies and this allows for a detailed analysis of a fundamental tradeoff between the intensity and duration of punishment. Second, our paper contributes to the literature on costly political conflict by providing a formal framework for investigating the transitional dynamics between conflict and cooperation.⁶ In particular, our model bears a similar structure to Yared (2009), though in contrast to this work, we introduce variable intervention intensity which allows for payoffs below the repeated static Nash equilibrium. This implies that, in contrast to this work, phases of intervention cannot last forever and must necessarily precede phases of cooperation. Third, our paper contributes to the literature on punishments dating back to the work of Becker (1968). In contrast to this work which considers static models, we consider a dynamic environment in which the government lacks the commitment to punish.⁷ This allows for an analysis of the optimal time structure of punishments together with the tradeoff between the duration and intensity of punishments.

The paper is organized as follows. Section 2 describes the model. Section 3 defines the efficient sequential equilibrium. Section 4 characterizes the equilibrium and provides our main results. Section 5 provides extensions and some discussion. Section 6 concludes. The Appendix contains all proofs and additional material not included in the text.

2 Model

We consider a dynamic environment in which a principal seeks to induce an agent into limiting disturbances. In every period, the principal has two options. On the one hand, he can forcefully intervene to control disturbances himself, and in doing so he chooses the intensity of force. On the other hand, the principal can withhold force and allow the agent to exert *unobservable* effort in controlling disturbances. In this situation, if a disturbance occurs, the principal cannot determine if it is due to the agent's negligence or due to bad luck. In addition to this informational asymmetry, both the principal and the agent suffer from limited commitment. In our benchmark model, we rule out payments from the principal to the agent—which are standard in the dynamic principal-agent literature—since

⁶Some examples of work in this literature are Acemoglu and Robinson (2006), Anderlini, Gerardi, and Lagunoff (2009), Baliga and Sjöström (2004), Chassang and Padró i Miquel (2009), Esteban and Ray (2008), Fearon (1995), Jackson and Morelli (2008), and Powell (1999). Schwarz and Sonin (2004) show that the ability to commit to randomizing between costly conflict and cooperation can induce cooperation. We do not assume the ability to commit to randomization, and the realization of costly conflict is driven by future expectations.

⁷Some examples of models of punishments are Acemoglu and Wolitzky (2009), Dal Bó and Di Tella (2003), Dal Bó, Dal Bó and di Tella (2006), and Polinski and Shavell (1979, 1984). We discuss our relationship to the literature on punishments in greater detail in Section 4.2.

our focus on is on the use of interventions. This is done purely for expositional simplicity. We allow for payments in Section 5.1 and show that none of our results are altered.

More formally, there are time periods $t = \{0, \dots, \infty\}$ where in every period t , the principal (p) and the agent (a) repeat the following interaction. The principal publicly chooses $f_t = \{0, 1\}$, where $f_t = 1$ represents a decision to intervene. If $f_t = 1$, then the principal publicly decides the intensity of force $i_t \geq 0$. In this case, the payoff to the principal is $-\pi_p \chi - Ai_t$ and the payoff to the agent is $w_a - g(i_t)$, where $A > 0$ and $g'(\cdot), -g''(\cdot) > 0$ with $g(0) = 0$, $g'(0) = \infty$, and $\lim_{i_t \rightarrow \infty} g'(i_t) = 0$. The concavity of $g(\cdot)$ captures the fact that there are diminishing returns to the use of intensity by the principal. The parameter A captures the marginal cost of intensive force.⁸ Within the term $-\pi_p \chi - Ai_t$ is embedded the cost of a stochastic disturbance, where π_p represents the probability of such a disturbance and χ represents its cost to the principal.⁹ Analogously, within the term $w_a - Ag(i_t)$ is the cost of the damage suffered by the agent.¹⁰

Importantly, conditional on intervention by the principal, both the principal and the agent are strictly better off under minimal force. Intuitively, choosing $i_t > 0$ imposes more physical damage on the agent. Moreover, it is statically inefficient from the perspective of the principal since it is more costly to use and does not diminish the likelihood of a disturbance. Therefore, conditional on $f_t = 1$, the principal would always choose $i_t = 0$ in a one-shot version of this game.

The principal can also decide to not intervene by choosing $f_t = 0$. In this case, the agent *privately* chooses whether to exert high effort ($e_t = \eta$) or low effort ($e_t = 0 < \eta$) in preventing a disturbance. Nature then stochastically chooses the realization of a publicly observed disturbance $s_t = \{0, 1\}$, where $s_t = 0$ represents the absence of a disturbance. If a disturbance does not occur, the principal receives 0, and if it occurs, the principal receives $-\chi$. Independently of the shock, the agent loses e_t from exerting effort. The stochastic realization of a disturbance occurs as follows. If $e_t = \eta$, then a disturbance occurs with probability $\pi_a(\eta) \in (0, 1)$ and if $e_t = 0$, then it occurs with probability $\pi_a(0) \in (\pi_a(\eta), 1]$. Therefore, high effort reduces the likelihood of a disturbance.¹¹ The

⁸For instance, A can decline if there is less international rebuke for the use of force.

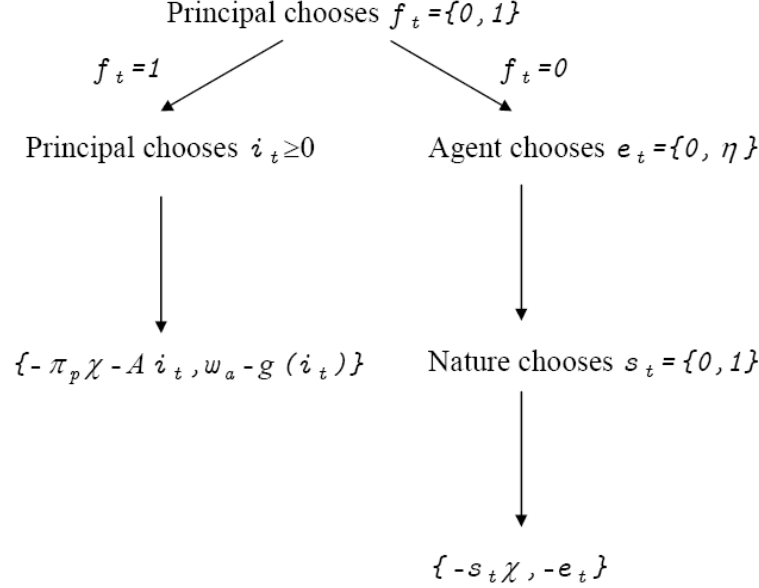
⁹That we have chosen the minimum level of intensity to be zero is only a normalization and has no effect on our results. More generally, one can interpret $i_t = 0$ as the statically optimal level of intensity under intervention.

¹⁰In practice, the agent can be a leader, a political party, or an entire society. In situations in which the agent is a group, the damage suffered by the agent can involve the killing of members of the group.

¹¹Due to the variety of applications, we do not take a stance on microfounding the source of disturbances. One can interpret these disturbances as being generated by a short-lived player who benefits from their realization (such as cross border raids into the Roman Empire by Germanic tribes) and who is less successful under intervention by the principal or high effort by the agent. Moreover, the realization of a disturbance could stochastically force the principal to make a permanent concession beneficial to this

parameter η captures the cost of effort to the agent.¹² The game is displayed in Figure 1.

Figure 1: Game



Let $u_j(f_t, i_t, e_t, s_t)$ represent the payoff to j at t , where value of i_t is only relevant if $f_t = 1$ and the values of e_t and s_t are only relevant if $f_t = 0$. Each player j has a period zero welfare

$$E_0 \sum_{t=0}^{\infty} \beta^t u_j(f_t, i_t, e_t, s_t), \beta \in (0, 1).$$

We make the following assumptions.

Assumption 1 (inefficiency of intervention) $\pi_p > \pi_a(\eta)$ and $-\eta > w_a$.

Assumption 2 (desirability of intervention) $\pi_a(0) > \pi_p$.

Assumption 1 states that, relative to payoffs under intervention, both the principal and the agent are strictly better off if the agent exerts high effort in preventing a disturbance. Intuitively, the agent is better informed about the sources of disturbances and is better than the principal at preventing them. Moreover, from an ex-ante perspective, the agent

player. Under this interpretation, the principal may be able to unilaterally make a concession to end all disturbances, a situation we consider in Section 5.2.

¹²The cost can rise for instance if it becomes more politically costly for the agent to antagonize rival factions contributing to the disturbances. Alternatively, the agent might actually have an increased preference for disturbances. In this case, without affecting any of our results, one can modify the interpretation so that e_t subsumes the fact that the agent receives utility from the realization of a disturbance.

prefers to exert high effort to prevent a disturbance versus enduring the damage from any intervention by the principal.

Assumption 2 states that the principal is strictly better off using intervention to prevent a disturbance versus letting the agent exert low effort in preventing such a disturbance. This assumption has an important implication. Specifically, in a one-shot version of this game, $f_t = 1$ and $i_t = 0$ is the unique static Nash equilibrium. This is because conditional on $f_t = 0$, the agent chooses $e_t = 0$. Thus, by Assumption 2, the principal chooses $f_t = 1$ and $i_t = 0$. Since the agent cannot commit to controlling disturbances, the principal must intervene to do so himself.¹³ We refer to this situation with $f_t = 1$ and $i_t = 0$ as *direct control*.

Permanent direct control is always a sequential equilibrium of the repeated game. However, since it is inefficient (by Assumption 1), one can imagine that repeated game strategies can enhance the welfare of both players. Nevertheless, there are three political economy frictions to consider. First, the principal cannot commit to refraining from using intervention in the future, since he also suffers from limited commitment. Moreover, he cannot commit to using more than minimal force under intervention. Second, the agent cannot commit to choosing high effort. Finally, the principal does not observe the effort by the agent. Consequently, if a disturbance occurs, the principal cannot determine if this is accidental (i.e., $e_t = \eta$) or if this is intentional (i.e., $e_t = 0$).

Note that our simple benchmark model ignores four additional issues. First, as we mentioned, it ignores the possibility that the principal can pay the agent for reducing disturbances. Second, it ignores the possibility that permanent concessions by the principal can reduce the presence of disturbances in the future. Third, it ignores the possibility that the agent's identity can change over time because of political transitions. Fourth, it ignores the possibility that high intensity levels by the principal today can raise the cost of effort by the agent in the future, for example if the agent becomes more antagonistic. These issues are discussed in Section 5 which shows that our main conclusions are unchanged.

¹³Assumption 2 facilitates exposition by guaranteeing a unique long run equilibrium. If it is violated, the worst punishment to the principal is redefined as equal to $-\pi_a(0)\chi/(1-\beta)$ and none of our main results are changed. Section 5.2 provides an extension with a permanent concession which is isomorphic to this scenario.

3 Equilibrium Definition

In this section, we present our recursive method for the characterization of the efficient sequential equilibria of the game. We provide a formal definition of these equilibria in the Appendix. The important feature of a sequential equilibrium is that each player dynamically chooses his best response given the strategy of his rival at every public history.¹⁴

Since we are concerned with optimal policy, we characterize the set of equilibria which maximize the period 0 welfare of the principal subject to providing the agent with some minimal period 0 welfare U_0 . The most important feature of these equilibria due to the original insight achieved by Abreu (1988) is that they are sustained by the worst punishment. More specifically, all public deviations from equilibrium actions by a given player lead to his worst punishment off the equilibrium path, which we denote by $\underline{J}$ for the principal and $\underline{U}$ for the agent. Note that

$$\begin{aligned}\underline{J} &= -\frac{\pi_p \chi}{1-\beta} \text{ and} \\ \underline{U} &\leq \frac{w_a}{1-\beta}\end{aligned}$$

since the principal cannot receive a lower payoff than under permanent direct control which he can always choose. Moreover, the agent can be credibly punished by the principal at least as harshly as under permanent direct control.

Note that in characterizing this equilibrium, we take into account that it may be efficient for players to choose correlated strategies so as to potentially randomize over the choice of intervention, intensity, and effort. Let $z_t = \{z_t^1, z_t^2\} \in Z \equiv [0, 1]^2$ represent a pair of i.i.d. publicly observed random variables independent of s_t , of all actions, and of each other, where these are drawn from a bivariate continuous c.d.f. $G(\cdot)$. Let z_t^1 be revealed prior to the choice of f_t so as to allow the principal to randomize over the use of intervention and let z_t^2 be revealed immediately following the choice of f_t so as to allow the principal to randomize over intensity or the agent to randomize over the effort.

As is the case in many incentive problems, an efficient sequential equilibrium can be represented in a recursive fashion, and this is a useful simplification for characterizing equilibrium dynamics.¹⁵ Specifically, at any public history, the entire public history of the game is subsumed in the continuation value to each player, and associated with these

¹⁴Because the principal's strategy is public by definition, any deviation by the agent to a non-public strategy is irrelevant (see Fudenberg, Levine, and Maskin, 1994).

¹⁵This is a consequence of the insights from the work of Abreu, Pearce, and Stacchetti (1990).

two continuation values is a continuation sequence of actions and continuation values. More specifically, let U represent the continuation value of the agent at a given history. Associated with U is $J(U)$, which represents the highest continuation value achievable by the principal in a sequential equilibrium conditional on the agent achieving a continuation value of U . More formally, letting $\delta = \{f_z, i_z, e_z, U_z^F, U_z^H, U_z^L\}_{z \in Z}$, the recursive program which characterizes the efficient sequential equilibrium is

$$J(U) = \max_{\delta} \int \left[\begin{array}{l} f_z (-\pi_p \chi - A i_z + \beta J(U_z^F)) + \\ (1 - f_z) (-\pi_a(e_z) \chi + \beta ((1 - \pi_a(e_z)) J(U_z^H) + \pi_a(e_z) J(U_z^L))) \end{array} \right] dG_z \quad (1)$$

s.t.

$$U = \int \left[\begin{array}{l} f_z (w_a - g(i_z) + \beta U_z^F) + \\ (1 - f_z) (-e_z + \beta ((1 - \pi_a(e_z)) U_z^H + \pi_a(e_z) U_z^L)) \end{array} \right] dG_z, \quad (2)$$

$$J(U_z^F), J(U_z^H), J(U_z^L) \geq \underline{J} \quad \forall z \quad (3)$$

$$U_z^F, U_z^H, U_z^L \geq \underline{U} \quad \forall z \quad (4)$$

$$-\pi_p \chi - A i_z + \beta J(U_z^F) \geq \underline{J} \quad \forall z \quad (5)$$

$$\beta (U_z^H - U_z^L) (\pi_a(0) - \pi_a(e_z)) \geq e_z \quad \forall z \quad (6)$$

$$f_z \in [0, 1], i_z \geq 0, \text{ and } e_z = \{0, \eta\} \quad \forall z. \quad (7)$$

(1) represents the continuation value to the principal written in a recursive fashion at a given history. f_z , i_z , and e_z represent the use of intervention, the choice of intensity, and the choice of effort, respectively, conditional on today's random public signal $z = \{z^1, z^2\}$. U_z^F represents the continuation value promised to the agent for tomorrow conditional on intervention being used today at z . If intervention is not used, then the continuation value promised to the agent for tomorrow conditional on z is U_z^H if $s = 0$ (there is no disturbance) and U_z^L if $s = 1$ (there is a disturbance). Note that f_z depends only on z^1 since it is chosen prior to the realization of z^2 , but all other variables depend on z^1 as well as z^2 .

Equation (2) represents the promise keeping constraint which ensures that the agent is achieving a continuation value of U . Constraints (7) ensure that the allocation is feasible. Constraints (3) – (6) represent the incentive compatibility constraints of this game. Without these constraints, the solution to the problem starting from an initial U_0 is simple: The principal refrains from intervention forever. Constraints (3) – (6) capture the inefficiencies introduced by the presence of limited commitment and imperfect information which ultimately lead to the need for intervention. Constraint (3) captures

the fact that at any history, the principal cannot commit to refraining permanent direct control which provides a continuation welfare of $\underline{J}$. Constraint (4) captures the fact that at any history, the agent cannot commit to high effort, as he can choose low effort forever and ensure himself a continuation value of at least $\underline{U}$. Importantly, constraint (5) captures the fact that at any history, the principal cannot commit to using intensive force since this is costly. Constraint (5) ensures that the principal prefers to use intensive force and be rewarded for it in the future compared to his best deviation which involves using intervention with zero intensive force forever. Constraints (3) – (5) capture the constraint of limited commitment. Under perfect information, they imply that if players are sufficiently patient, the permanent absence of intervention can be sustained by the off-equilibrium threat of intervention. Constraint (6) captures the additional constraint of imperfect information: If the principal requests $e_z = \eta$, the agent can always privately choose $e_z = 0$ without detection. Constraint (6) ensures that the agent’s punishment from this deviation is weakly exceeded by the equilibrium path reward for choosing high effort.¹⁶

4 Analysis

We focus our analysis on the likelihood, the intensity, and the duration of intervention which are formally defined below.

Definition 1 (i) *The likelihood of intervention at t is $\Pr \{f_{t+1} = 1 | f_t = 0 \text{ and } s_t = 1\}$,* (ii) *the intensity of intervention at t is $E \{i_t | f_t = 1\}$, and (iii) the duration of intervention at t is $\Pr \{f_{t+1} = 1 | f_t = 1\}$.*

This definition states that the likelihood of intervention is the probability that the principal intervenes following a disturbance; the intensity of intervention is the expected intensity of the force used by the principal; and the duration of the intervention is the probability that intervention continues into the next period.

We also focus on *long run* equilibrium dynamics. We do so because these dynamics can be explicitly characterized in closed form, and because we can show that phases of intervention occur only in the long run.¹⁷ More specifically, we first show in Section 4.1 that the optimal contract in the long run is characterized by two phases of play: a cooperative phase and a punishment phase, where these two phases sustain each other. Second, we

¹⁶Note that we have ignored the constraint that the agent does not deviate to high effort if $e_z = 0$ since such a constraint never binds in equilibrium.

¹⁷See Yared (2009) for a similar model which more explicitly describes short run transitional dynamics.

describe in Section 4.2 an important tradeoff in the optimal contract between the duration and the intensity of intervention. Finally, in Section 4.3 we consider comparative statics.

To facilitate exposition, we assume that players are sufficiently patient for the remainder of our discussion.

Assumption 3 (High Patience) $\beta > \hat{\beta}$.

The exact value of $\hat{\beta}$ is described in the Appendix.¹⁸

4.1 Characterization

Let

$$\delta^*(U) = \{f_z^*(U), i_z^*(U), e_z^*(U), U_z^{F*}(U), U_z^{H*}(U), U_z^{L*}(U)\}_{z \in Z}$$

represent an argument which solves (1) – (7). Since $\delta^*(U)$ may not be unique, we focus on the unique solution which satisfies the *Bang-Bang* property as described by Abreu, Pearce, and Stacchetti (1990).¹⁹ In our context, the *Bang-Bang* property is satisfied if the equilibrium continuation value pairs at t following the realization of z_t^1 are extreme points in the set of sequential equilibrium continuation values. Define

$$\bar{U} = -\frac{\pi_a(0)\eta}{(1-\beta)(\pi_a(0) - \pi_a(\eta))}. \quad (8)$$

Let $\lim_{t \rightarrow \infty} \Pr \{U_t \leq \bar{U}\}$ represent the long run probability that the agent receives a continuation value (following the realization of z_t^1) which is weakly below $\bar{U}$ in the solution to the program.

Proposition 1 (*characterization*)

1. $\lim_{t \rightarrow \infty} \Pr \{U_t \leq \bar{U}\} = 1 \ \forall U_0$, and

¹⁸This assumption guarantees that the likelihood of punishment is bounded away from 1 and that the duration of punishment is bounded away from 0, which guarantees that the long run equilibrium can be explicitly characterized. The value of $\hat{\beta}$ is below 1 as long as η is sufficiently bounded away from w_a so that permanent reversion to the static Nash equilibrium is a sufficient enough threat to induce high effort.

¹⁹Efficient equilibria which do not satisfy the *Bang-Bang* property emerge here in part because information is coarse, an issue which is discussed in Yared (2009). The *Bang-Bang* equilibrium we describe is the unique optimum if players are constrained to one-period memory and if a rich and asymptotically uninformative public signal of the agent's effort were available to the principal. Details available upon request.

2. If $U \leq \bar{U}$, then $E f_z^*(U) = (\bar{U} - U) / (\bar{U} - \underline{U})$ and $\forall z$

$$\begin{aligned} i_z^*(U) &= i^*, \\ e_z^*(U) &= \eta, \\ U_z^{F^*}(U) &= (\underline{U} - w_a + g(i^*)) / \beta, \\ U_z^{H^*}(U) &= \bar{U}, \text{ and} \\ U_z^{L^*}(U) &= \bar{U} - \eta / (\beta (\pi_a(0) - \pi_a(\eta))) \end{aligned}$$

for i^* and $\underline{U}$ which satisfy

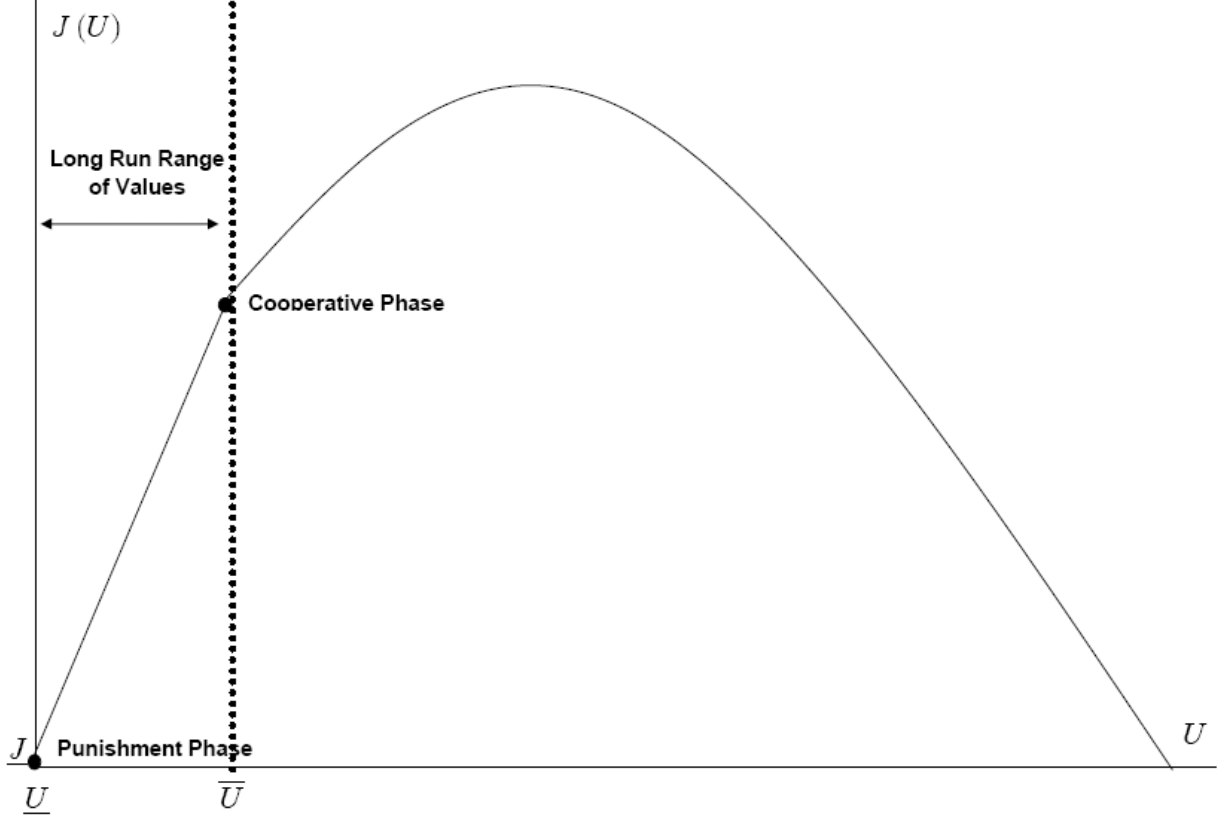
$$\begin{aligned} 1 &= \frac{g'(i^*) (\pi_p - \pi_a(\eta)) \chi + A i^*}{A (-\eta - w_a + g(i^*))}, \text{ and} \\ \underline{U} &= -\frac{(\pi_p - \pi_a(\eta)) \chi \eta + A i^* (w_a - g(i^*))}{(1 - \beta) ((\pi_p - \pi_a(\eta)) \chi + A i^*)}. \end{aligned} \tag{9}$$

This proposition states that in the long run, continuation values are weakly below $\bar{U}$ and it explicitly characterizes the solution for $U \leq \bar{U}$. More specifically, in the long run, the principal exerts a unique level of intensity i^* , the agent exerts high effort, and continuation values for tomorrow are conditioned on whether or not intervention is used and whether or not a disturbance occurs in the absence of intervention. The continuation value U is therefore provided to the agent by randomizing over a *cooperative phase* and a *punishment phase*. In the cooperative phase, intervention is not used and the agent and principal receive $\bar{U}$ and $J(\bar{U})$, respectively, following the realization of z_t^1 . In the punishment phase, intervention is used and the agent and principal receive $\underline{U}$ and $J(\underline{U}) = \underline{J}$, respectively, following the realization of z_t^1 .

More specifically, in the cooperative phase at t , the principal does not intervene ($f_t = 0$) and the agent chooses high effort ($e_t = \eta$). If there is no disturbance at t ($s_t = 0$), then the cooperative phase at $t + 1$ occurs with probability 1. If there is a disturbance at t ($s_t = 1$), then the cooperative phase at $t + 1$ occurs with probability $1 - l^*$, and the punishment phase at $t + 1$ occurs with probability l^* . In contrast, in the punishment phase at t , principal chooses intervention ($f_t = 1$) and a unique level of intensity $i_t = i^*$. The punishment phase at $t + 1$ occurs with probability d^* and the cooperative phase at $t + 1$ occurs with probability $1 - d^*$. Note that given Definition 1, it is clear from this characterization that the optimal likelihood, intensity, and duration of intervention correspond to l^* , i^* , and d^* , respectively, and these can be characterized explicitly in our

framework.²⁰

Figure 2: $J(U)$



To understand the first part of Proposition 1 consider Figure 2 which depicts $J(U)$ as a function of U . The y -axis represents $J(U)$ and the x -axis represents U , with $\bar{U}$ situated on the x -axis. Note that an efficient equilibrium necessarily begins on the downward sloping portion of $J(U)$ since it is not possible to make the principal better off along this portion without making the agent worse off. Along the upward sloping portion of $J(\cdot)$, both the principal and agent can be made better off from an increase in U since this is associated with a lower probability realization of intervention which is costly to both players. Along the downward sloping portion of $J(\cdot)$, the principal is made worse off from an increase in U since this is associated with a higher probability realization of low effort by the agent which is costly to the principal but beneficial to the agent. Along the equilibrium path, whenever the principal requests high effort from the agent, he rewards the agent for the absence (realization) of a disturbance with an increase (decrease) in continuation value. Therefore, a sequence of disturbances can cause the continuation value to the agent to

²⁰More specifically, $U_z^{F*}(\underline{U}) = (1 - d^*)\bar{U} + d^*\underline{U}$ and $U_z^{L*}(\bar{U}) = (1 - l^*)\bar{U} + l^*\underline{U}$.

decline below $\bar{U}$.²¹

$\bar{U}$ is important for two reasons. First, it can be shown that if $U \geq \bar{U}$, $f_z^*(U) = 0 \forall z$ so that intervention is used with zero probability. The reason is that it is too costly for both the principal and for the agent and hence it is inefficient to use it when the promised value U is not low. Now suppose there was zero probability of continuation values traveling below $\bar{U}$. Then there would be zero probability of intervention along the equilibrium path, and the agent would optimally choose low effort forever. This would obviously violate the incentive compatibility constraint of the principal by Assumption 2.²² Therefore, intervention must occur along the equilibrium path to induce high effort and continuation values must eventually decline below $\bar{U}$. Second, $\bar{U}$ is important because continuation values in the future cannot increase above $\bar{U}$ once they have declined below $\bar{U}$. The intuition is that conditional on being forced to use high intensity interventions with some frequency, the principal wants to extract as much as possible from the agent in the periods when he does not intervene. He does this by always requesting high effort, which hurts the agent and benefits the principal. If instead continuation values were to increase above $\bar{U}$, this would imply that the agent would be able to exert low effort at some point in the future. However, this would force the principal must to punish the agent more often or more intensely in order to satisfy (2), which is costly. In other words, the least costly way of implementing credible punishments is an equilibrium in which once punishments are used the agent is never again allowed to exert low effort. Allowing the agent to exert low effort in the future (and thus moving above $\bar{U}$) only serves to make cooperation less acceptable to the principal, which makes him less willing to punish with the same intensity.²³

The intuition behind the second part of Proposition 1 is that in equilibrium, phases of cooperation and phases of punishment sustain each other. In the cooperative phase, the agent exerts high effort because he knows that failure to do so raises the probability of a disturbance which can trigger a transition to the punishment phase. In the punishment phase, the principal temporarily intervenes with a unique level of intensive force. The principal exerts costly force since he knows that failure to do so would trigger the agent to choose low effort in all future cooperative phases, making direct control—i.e., permanent

²¹For more details, see the Appendix.

²²In a model which allows for payments from the principal to the agent, the second part of Proposition 1 holds exactly, though the first part may not necessarily do so since a long enough absence of disturbances can lead to the permanent absence of intervention. See Section 5.1 for a discussion.

²³Technically, if $U_z^{H*}(U) > \bar{U}$, then (6) would not bind which is inefficient by the concavity of $J(\cdot)$.

Note that the first part of Proposition 1 holds for all solutions, not just those which satisfy the *Bang-Bang* property.

intervention with minimal intensity—a necessity.

Importantly, the values of $\underline{U}$ and $J(\bar{U})$ are intimately linked. To see why, consider the system of equations which characterizes the long run equilibrium:

$$\bar{U} = -\eta + \beta \left((1 - \pi_a(\eta) l^*) \bar{U} + \pi_a(\eta) l^* \underline{U} \right) \quad (10)$$

$$\underline{U} = w_a - g(i^*) + \beta \left(d^* \underline{U} + (1 - d^*) \bar{U} \right) \quad (11)$$

$$J(\bar{U}) = -\pi_a(\eta) \chi + \beta \left((1 - \pi_a(\eta) l^*) J(\bar{U}) + \pi_a(\eta) l^* \underline{J} \right) \quad (12)$$

$$\underline{J} = -\pi_p \chi - A i^* + \beta \left((1 - d^*) J(\bar{U}) + d^* \underline{J} \right). \quad (13)$$

(10) and (11) represent the continuation value to the agent during cooperation and punishment, respectively. (10) shows that in the cooperative phase, the agent exerts high effort today and faces two possibilities tomorrow. If a disturbance occurs and he is not forgiven, play moves to punishment and he obtains $\underline{U}$. Otherwise, cooperation is maintained and he receives $\bar{U}$ tomorrow. (11) shows that in the punishment phase, the agent endures punishment with intensity i^* today, and he receives $\underline{U}$ tomorrow with probability d^* and $\bar{U}$ tomorrow with probability $1 - d^*$. (12) and (13) are analogously defined for the principal. In particular, (12) shows that during cooperation the principal suffers from disturbances with probability $\pi_a(\eta)$, and (13) shows that during punishment the principal suffers from disturbances with a higher probability π_p and he also suffers from intervening with force.²⁴

Crucially, the value of $\bar{U}$ does not depend on the value of i^* since $\bar{U}$ is self-generating in equilibrium.²⁵ Moreover, as discussed in Section 3, $\underline{J}$ is independent of i^* because it simply corresponds to the repeated static Nash payoff to the principal—i.e., direct control. Therefore, (10) – (13) is a system of four equations and five unknowns— $\underline{U}$, $J(\bar{U})$, l^* , i^* , and d^* —where the value of i^* is selected to maximize $J(\bar{U})$.

Note that this system of equations allows us to trace exactly how the cooperative and punishment phases sustain each other. Since $\bar{U}$ is exogenously determined, equation (10) implies that the lower is $\underline{U}$, then the lower is the implied value of l^* . Intuitively, the harsher the punishment, the less often it needs to be used. Because $\underline{J}$ is also exogenous, (12) shows that $J(\bar{U})$ is decreasing in l^* . Since payoffs under intervention are fixed for the principal, he is better off if he needs to intervene less often. As a consequence, the highest possible $J(\bar{U})$ is attained by the lowest $\underline{U}$, as this makes for the longest sustainable cooperative phase.

²⁴Note that equations (10) and (13) naturally emerge from equations (6) and (5), the incentive compatibility constraints on the agent and principal, respectively.

²⁵That is, $\bar{U}$ is derived by combining (2) with (6) (which binds) given that $e_z^*(\bar{U}) = \eta$ and $U_z^{H^*}(\bar{U}) = \bar{U}$.

Similarly, equations (11) and (13) imply that, conditional on i^* , the higher is $J(\bar{U})$, then the higher is the implied value of d^* , and the lower is the implied value of $\underline{U}$. This is because the higher the principal's welfare under cooperation, the more easily can the principal be induced to punish for longer. Longer punishments lower $\underline{U}$ which again increases $J(\bar{U})$. Consequently, the optimal i^* that maximizes the principal's value of cooperation simultaneously also minimizes the agent's value of punishment.

4.2 Tradeoff between Intensity and Duration of Intervention

In this section, we consider the choice of intensity in the optimal contract together with its implications for the optimal likelihood and duration of intervention. In doing so, we highlight a fundamental tradeoff between the intensity and duration of intervention.

To this end, it is useful to consider the wider implications of the system given by (10) – (13). In particular, consider an exogenous level of intensity i —i.e., not necessarily the optimal level i^* . For a given i , this system of equations is linear in four unknowns and it is therefore solvable. Take the solutions for l^* and d^* given i , and call them $l(i)$ and $d(i)$ as they are now a function of the exogenous level of i that we are considering. In other words, $l(i)$ corresponds to the likelihood of intervention under intensity i and $d(i)$ corresponds to the duration of intervention under intensity i .

Proposition 2 (optimal intervention) *The optimal levels of l^* , i^* , and d^* satisfy $l^* = l(i^*)$ and $d^* = d(i^*)$ for i^* defined in (9) where $l(\cdot)$ and $d(\cdot)$ are continuously differentiable functions with $l'(i) < (>) 0$ if $i < (>) i^*$ and $d'(i) < 0$.*

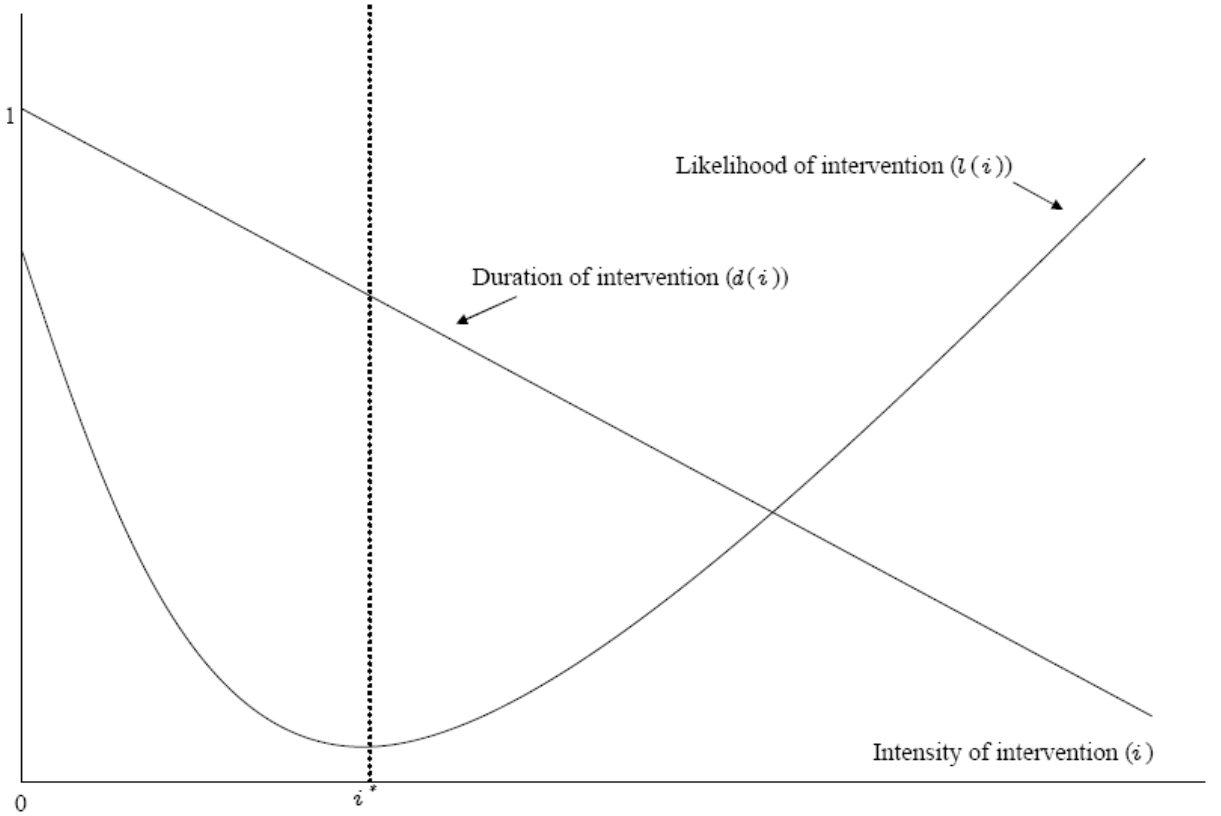
Proposition 2 states that, in the set of equilibria with the same structure as the efficient equilibrium, an increase in intensity reduces the likelihood of intervention for $i < i^*$ and it increases the likelihood of intervention for $i > i^*$. Moreover, an increase in intensity always reduces the duration of intervention. This proposition implies that there is a tradeoff between the intensity and duration of intervention, and that the optimal level of intensity i^* corresponds to the point which minimizes the likelihood of intervention. This proposition is displayed graphically in Figure 3, where intensity i is on the x -axis and the implied likelihood and duration of intervention— $l(i)$ and $d(i)$, respectively—are on the y -axis.

The principal's incentives to intervene are the driving force behind Proposition 2. Again, recall that the principal can always deviate to permanent direct control, which gives him a fixed exogenous payoff. As a consequence, if the intensity of intervention rises, then the principal can only be induced to exert this level of intensity if the resumption

of cooperation following intervention is more likely. This is the logic behind (13) and it implies that $d'(i) < 0$, so that the duration of intervention is declining in intensity.

Now consider what this implies for the welfare of the agent under punishment, $\underline{U}$. At low levels of i , an increase in intensity naturally means that the prospect of punishment is worse for the agent, and $\underline{U}$ decreases in i . However, at higher levels of i , diminishing returns set in and the smaller marginal increase in pain $g'(i)$ is outweighed by the reduction in punishment duration implied by (13). As a consequence, above a certain i , $\underline{U}$ becomes *increasing* in i .

Figure 3: Likelihood, Intensity, and Duration of Intervention



Since the agent's value under punishment first decreases and then increases with intensity, the likelihood of intervention $l(i)$ first decreases and then increases with intensity, as implied by (10). As the punishment for the agent becomes worse, a smaller likelihood of punishment is needed to satisfy (10). As discussed above, lower likelihood is better from the perspective of the principal because it maximizes the duration of cooperation. Therefore, the principal always chooses the level of intensity that minimizes likelihood. As stated in Proposition 2, this level is i^* .

As an aside, note that our selection of an interior point i^* relies on our assumption that $g'(0)$ is sufficiently high. If $g'(0)$ were small, then one could construct environments in which $i^* = 0$ so that indirect control is not sustainable and the principal resorts to permanent direct control, as in Yared (2009). Intuitively, the punishment to the agent is not sufficiently dire to warrant its use by the principal. Moreover, note that the uniqueness of i^* defined in (9) is guaranteed by the global concavity of $g(\cdot)$. If instead $g(\cdot)$ were weakly convex, there would be no tradeoff between the duration and intensity of intervention, and the optimal level of intensity would be either 0 or the maximal feasible level of intensity.

These results are related to static models of punishment which study a variety of situations, such as extortion and slavery.²⁶ They are also related to the law and economics literature which considers the tradeoff between the likelihood of punishment (i.e., the probability of capturing criminals) and the harshness of punishment (i.e., the length of incarceration).²⁷ As in our environment, this literature establishes that choosing the harshest existing punishment is suboptimal because costly punishments must be exercised *in equilibrium*. Second, the law and economics literature highlights a complementarity between the likelihood and the harshness of punishment which is also present in our framework. More specifically, in our model an increase l^* and a reduction in $\underline{U}$ are complementary tools for the reduction of the punishment continuation value $U^L(\bar{U})$. Nonetheless, in contrast to our dynamic model, static models by definition cannot distinguish between the intensity and the duration of punishment, and hence they cannot provide any answers to the motivating questions of our analysis. In this regard, the tradeoff in our model between the intensity and duration of punishment and its relationship to the absence of commitment on the side of the principal is novel to the literature on punishment.²⁸

4.3 Comparative Statics

In this section, we consider the effect of three factors on the optimal likelihood, intensity, and duration of intervention. First, we consider the effect of a decline in the cost of intensity to the principal (A). Second, we consider the effect of a rise in the cost of disturbances to the principal (χ).²⁹ Finally, we consider the effect of a rise in the cost of effort to the agent (η), where this can occur for instance if it becomes more politically costly

²⁶See Dal Bó and Di Tella (2003) and Dal Bó, Dal Bó and di Tella (2006) for an application to political capture and Chwe (1990) and Acemoglu and Wolitzky (2009) for labor contracts with limited liability.

²⁷See, for instance, the seminal articles by Becker (1968) and Polinsky and Shavell (1979,1984).

²⁸Because applying punishments is costly to the principal, static models need to assume that the principal can commit to some punishment intensity as a function of observable outcomes.

²⁹One can also interpret this parameter as reflecting the preferences of the principal, so an increase in χ reflects a transition to a principal who is less tolerant of disturbances.

for the agent to antagonize rival factions contributing to disturbances or alternatively if he acquires a higher preference for the realization of disturbances. We make the following assumption to facilitate our discussion.

Assumption 4 $g(i) = i^\theta$ for $0 < \theta < 1$.

As we discuss further below, the only purpose of this assumption is to make the effect on duration of a change in A or χ unambiguous. The comparative statics are summarized in the below proposition.³⁰

Proposition 3 (*comparative statics*)

1. If A decreases (increases), then l^* decreases (increases), i^* increases (decreases), and d^* decreases (increases),
2. If χ increases (decreases), then l^* decreases (increases), i^* increases (decreases), and d^* decreases (increases), and
3. If η increases (decreases), then l^* increases (decreases), i^* increases (decreases), and d^* decreases (increases).

This proposition states that all three changes increase the optimal intensity and decrease the optimal duration of intervention. However, only the third change also raises its likelihood whereas the first two changes decrease its likelihood.

To see why intensity must rise, consider the first case. If the cost of intensity declines, then the principal's return to intensity rises since it is cheaper to provide incentives to the agent via intensive force.³¹ In the second case, if the cost of disturbances rise, the principal should raise the intensity of intervention since the return to delegating to the agent rises relatively to direct control. As direct control worsens, higher intensity becomes incentive compatible. In the third case, if the cost of effort for the agent rises, then it is harder for the principal to provide incentives to the agent with lower levels of intensity,

³⁰Performing comparative statics with respect to the probability of a disturbance is not straightforward given that this would affect the values of π_p , $\pi_a(\eta)$, and $\pi_a(0)$ jointly.

³¹This is arguably the case in some of our motivating examples, since international rebuke against the use of violence in restive regions changes over time and often causes governments to change their intervention strategy.

and higher levels of intensity become optimal.³² In all three cases, because the principal needs more inducement to use more intensive punishments, these increases in the level of intensity necessitate a decline in the duration of intervention.

Though all three changes increase the optimal intensity and decrease the optimal duration of intervention, only the third also raises its likelihood. Specifically, if the cost of intensity to the principal declines or if the cost of disturbances to the principal rises, then higher intensity slackens the agent's incentive constraint. As a consequence, the principal can afford to forgive him more often without weakening incentives, and the likelihood of intervention declines. In contrast, if the cost of effort to the agent rises, then incentives are harder to provide for the agent, so that likelihood of intervention must rise following the realization of a disturbance.

Note that the comparative statics with respect to the likelihood and the duration of intervention rely on the fact that the principal responds optimally to changes in the environment by increasing the level of intensity. To see why, consider the effect of each of these factors absent any change in the level of intensity, where the ensuing hypothetical suboptimal equilibrium can be constructed as in Section 4.2. Consider the effect of a decrease in the cost of intensity to the principal or an increase in the cost of disturbances to the principal absent any change in i . In this circumstance, the implied likelihood of intervention declines and implied duration of intervention *rises*. This is because it becomes easier to provide incentives to the principal to use force (i.e., either the cost of force is lower or the marginal benefit of resuming cooperation rises). Since incentives to the principal are easier to provide but i is fixed, the duration of intervention can rise. Therefore punishment becomes more severe for the agent, and the likelihood of intervention declines.³³ In contrast, when i^* is allowed to adjust, Proposition 3 shows that the increase in intensity is so large that it requires a *reduction* in the duration of intervention. This final comparative static relies on Assumption 4, and one can construct environments in which a decline in A or a rise in χ would barely change i^* , thereby

³²This comparative static is particularly fitting for understanding the case of the Roman Empire, which utilized more brutal force in the western region of the Empire—where chieftain control was tenuous and therefore needed higher effort—relative to the eastern regions—where client rulers had more control. Specifically, Luttwak (1975) writes:

"[T]he client rulers of the east normally enjoyed secure political control over their subjects...By contrast, in the less structured polities of Europe, the prudence of the well-informed would not necessarily restrain all those capable of acting against Roman interest...[O]ne can therefore say that while Roman military power was freely converted into political power vis-à-vis the sophisticated polities of the East, when employed against the primitive peoples of Europe its main use was the direct application of force." (p.32-33)

³³Formally, this is equivalent to stating that $d(i)$ is decreasing in A and increasing in χ .

generating an increase in the duration of intervention.³⁴

Analogously, one can consider the effect of a rise in the cost of effort to the agent, absent any change in i . In this circumstance, the implied likelihood of intervention rises and the implied duration of punishment declines. This is because it becomes more difficult to provide incentives to the agent to exert high effort, so that the likelihood of intervention rises, reducing the value of cooperation for the principal. Because the principal puts lower value on cooperation, the duration of intervention must decline so as to provide the principal with enough inducement to exert the same level of intensity. In this circumstance, the optimal level of intensity rises and therefore *mitigates* the rise in the likelihood of intervention, and this reinforces the decline in the duration of intervention.³⁵

5 Extensions and Discussion

Our benchmark model ignores four additional issues. First, it ignores the possibility that the principal can pay the agent for reducing disturbances. Second, it ignores the possibility that permanent concessions by the principal can reduce the presence of disturbances in the future. Third, it ignores the possibility that the agent's identity can change over time because of political transitions. Fourth, it ignores the possibility that high intensity levels by the principal today can raise the cost of effort by the agent in the future, for example if the agent becomes more antagonistic. These issues are discussed in the below four extensions which show that our main conclusions are unchanged.³⁶ Following these extensions, we discuss implications of our model for optimal counterinsurgency policy.

5.1 Temporary Payments

Our benchmark model ignores the presence of payments from the principal to the agent which are standard in principal-agent relationships. Consider an extension of our model where if the principal does intervene at t ($f_t = 0$), he chooses a payment $c_t \geq 0$ which he makes to the agent prior to the choice of effort by the agent. Thus, conditional on $f_t = 0$, the payoff to the principal at t is $-c_t - s_t\chi$ and the payoff to the agent is $c_t - e_t$. Under this extension, our model is isomorphic to Fong and Li (2009) with the exception that

³⁴This would be true for instance if $g(\cdot)$ features high curvature around i^* , for instance if $-i^*g''(i^*)/g'(i^*) > 1$.

³⁵The rise in the likelihood of intervention occurs independently of Assumption 4 since the principal must be strictly worse off if η rises.

³⁶Due to space restrictions, we describe these results informally, but more details are available upon request.

their model is a special case of ours with i_t constrained to 0 at every date.

Under this extension, the prospect of future payment can serve as a reward for the successful avoidance of disturbances and the use of intervention continues to serve as a punishment for disturbances. Moreover, payment is never used during intervention since the principal would like to make the agent suffer as much as possible. As such, the second part of Proposition 1, Proposition 2, and Proposition 3 are preserved.

More specifically, if a sufficient number of disturbances occur, then continuation values must decline below $\bar{U}$ defined in (8) and punishment necessarily occurs. Intuitively, because of limited liability, it is inefficient to provide incentives using payments alone, and it is efficient to use punishments in the form of intervention. Moreover, by analogous reasoning as in Proposition 1, continuation values cannot rise above $\bar{U}$ once they have declined below it. Therefore, continuation values must be trapped below $\bar{U}$ if intervention is ever used along the equilibrium path, and no payment will ever be made going forward in this situation.

The main difference between the benchmark model and the extended model is that under some conditions, the extended model admits another long run equilibrium in which intervention is not used.³⁷ In this alternate long run equilibrium which is described in Fong and Li (2009), the principal does not use intervention, and he only uses payment in the provision of incentives. More specifically, the long run equilibrium features a payment phase in which the principal pays the agent and a no-payment phase in which the principal does not pay the agent. In both phases, the principal requests high effort from the agent. The absence of a disturbance leads to a probabilistic exit from the no-payment phase and the presence of a disturbance leads to a probabilistic exit from the payment phase.

Thus, the equilibrium of the extended model can feature history-dependence in the long run contract. On the one hand, sufficient absences of disturbances can lead to an equilibrium which features no intervention and repeated payment.³⁸ On the other hand, a sufficient realization of disturbances can lead to an equilibrium which features no payment and repeated intervention as in our benchmark model.

5.2 Permanent Concessions

Consider an extension of our benchmark model where if the principal does not intervene at t ($f_t = 0$), he can choose a permanent concession which, with some abuse of notation,

³⁷This requires a condition which guarantees the existence of a trigger-strategy equilibrium in which payment induces high effort. Absent this condition, the unique long run equilibrium involves repeated intervention.

³⁸This is also the case if the initial condition U_0 is chosen to be sufficiently high.

we refer to as $c_t = \{0, 1\}$. If $c_t = 0$, then no concession is made and the rest of the period proceeds as in our benchmark model. In contrast, if $c_t = 1$, a permanent concession is made which ends the game and provides a continuation value J^C to the principal and U^C to the agent starting from t . Such a concession can come in the form of independence, land, or political representation, for instance, and we assume that it satisfies the agent and ends all disturbances. Specifically, suppose that $U^C > 0$, so that it provides the agent with more utility than low effort forever.

Clearly, if $J^C < \underline{J}$, then the principal cannot possibly be induced to make a concession since he prefers permanent direct control. Therefore, the equilibrium would be exactly as the one we have characterized. Conversely, if $J^C > -\pi_a(\eta)\chi/(1-\beta)$, then the efficient equilibrium involves no intervention since the concession provides a better payoff to the principal than the best payoff under indirect control. In this case, the principal simply makes the concession in period 0 and the game ends. We therefore consider the more interesting case in which $J^C \in (\underline{J}, -\pi_a(\eta)\chi/(1-\beta))$.

In this situation, the provision of this concession serves as a reward for the successful avoidance of disturbances and the use of intervention continues to serve as a punishment for disturbances.³⁹ Clearly, if a sufficient number of disturbances are avoided, then intervention never takes place and the long run equilibrium features the concession by the principal together with the end of all conflict so as to reward the agent for good behavior. In contrast, if a sufficient number of disturbances occur, then continuation values decline below $\bar{U}$ defined in (8) and punishment necessarily occurs. Moreover, by analogous reasoning as in Proposition 1, continuation values cannot rise above $\bar{U}$ once they have declined below it. Therefore, continuation values must be trapped below $\bar{U}$ if intervention is ever used along the equilibrium path, and no concession will ever be made going forward in this situation.

The equilibrium of the extended model thus admits two potential long run outcomes, one with a permanent concession and the other which is analogous in structure to the one which we consider. Thus, as in our benchmark environment, the second equilibrium features phases of cooperation and punishment which sustain each other, it features a tradeoff between the intensity and duration of intervention, and it features the same comparative statics. Nevertheless, the equilibrium is not quantitatively identical to the one in the benchmark model precisely because the min-max for the principal is now J^C as opposed to $\underline{J}$. In other words, the principal cannot experience a continuation value below that which he can guarantee himself by making a concession to the agent. This

³⁹This is because rewarding the agent by allowing low effort is inefficient for the principal as well as the agent.

implies that the agent's welfare under punishment $\underline{U}$ must be higher in the extended model. Thus, the likelihood of punishment is higher and its duration shorter because it is harder to provide incentives to the principal and to the agent.⁴⁰

As an aside, note that if the principal lacks commitment to concessions and if a concession costs the principal $J^C(1 - \beta)$ in every period, then nothing changes as long as $J^C > \underline{J}$, since concessions can be enforced. If instead $J^C < \underline{J}$, then temporary concessions may be featured along the equilibrium path, but the long run characterization of the equilibrium is exactly as in our benchmark model.

5.3 Political Transitions

Our model additionally ignores the role of political transitions since it assumes that the two players interact with each other forever. This issue is particularly relevant for the case of the agent since the dynamics of the equilibrium are generated by the need for the principal to punish the agent for the realization of past disturbances. Clearly, there is no need for the principal to punish an agent who cannot possibly be blamed for past disturbances.

To explore this issue further, imagine if in every period there is a probability $1 - q$ that the incumbent agent is exogenously replaced by another identical agent, where replacement yields an exogenous continuation value to the incumbent. Moreover, to simplify discussion, consider the efficient sequential equilibrium which maximizes the principal's period 0 welfare, where the optimal contract now clearly specifies the identity of the agent whom the principal faces.

It is easy to show that in such a setting, the second part of Proposition 1 will hold for the long run interaction between the principal and a given agent, where β in Proposition 1 and in (8) is replaced by βq which corresponds to the relevant discount factor for the agent.⁴¹ In other words, our characterization of the cooperative and punishment phases holds for the interaction between the principal and an agent after several disturbances have occurred during the agent's tenure. This equilibrium features phases of cooperation and punishment which sustain each other. Moreover, one can show that for q sufficiently close to 1, it features the same tradeoff between the intensity and duration of intervention, and it features same exact comparative statics. Nonetheless, the model is not quantitatively equivalent to our benchmark environment since the principal's and the agent's discount

⁴⁰We have implicitly assumed that an analogous condition to Assumption 3 holds so that the implied duration of punishment is bounded away from zero.

⁴¹This statement refers to the continuation value to the agent adjusted by the continuation value associated with replacement.

factors differ from one another. Moreover, one can show that as q declines, it becomes more difficult for the principal to provide incentives to the agent so that the likelihood of intervention rises and the duration of intervention declines.

An important new feature of the extended model which is not present in the benchmark model is that a political transition causes the continuation value to the agent to rise above $\bar{U}$. This is because it is inefficient for the principal to punish an agent who is not responsible for the exertion of effort in the past by providing him with low welfare. Note further that it is straightforward to combine this extension with that of Section 5.2 which allows the principal to make a permanent concession. In such a setting, the long run will always feature a permanent concession by the principal. This is because even if one agent is punished and may never receive the concession himself, there is always a positive probability going forward that the agent which replaces him will be successful at preventing disturbances and will therefore be rewarded with a permanent concession.

An additional issue to consider is the possibility that the principal can endogenously replace the agent with another identical agent via assassination or demotion. In this environment, we can ignore without any loss of generality the principal's incentives to replace an incumbent since this does not provide any additional welfare to the principal given that future agents are identical to the incumbent.⁴² Moreover, one can show that the optimal contract specifies either intervention or replacement as the optimal form of punishment. For example, if replacement is not sufficiently costly to the incumbent agent, then replacement is never used in equilibrium, and all of the results from our benchmark model hold. Alternatively, if replacement strictly dominates intervention as a form of punishment, then intervention is never used in equilibrium, and our model becomes analogous to the classical Ferejohn (1986) model of electoral control, with the exception that we consider history-dependent strategies.

5.4 Endogenous Effort Cost

Our model additionally ignores the fact that the use of intensity by the principal can potentially make it more difficult for the agent to exert effort in preventing disturbances. This would occur if the agent loses political credibility with the population he is supposed to control. To explore this issue further, imagine if the cost of high effort η depends on time so that it is denoted by η_t and it can either be low ($\eta_t = \eta^L$) or high ($\eta_t = \eta^H$).

⁴²Specifically, any out of equilibrium removal of an incumbent can prompt all future agents to punish the principal by exerting zero effort forever.

Suppose $\eta_0 = \eta^L$ and imagine the following process for η_t :

$$\eta_t = \begin{cases} \eta^H & \text{if } f_k = 1 \text{ and } i_k > \tilde{i} \text{ for any } k < t \\ \eta^L & \text{otherwise} \end{cases}.$$

This means that if the principal ever exceeds a certain level of intensity, then the cost of high effort for the agent permanently rises. Moreover, suppose $\tilde{i}$ is below the optimal level of intensity in an environment in which $\eta_t = \eta^L$ for all t . This means that if the principal uses the same level of intensity as in our benchmark environment, the cost of effort for the agent permanently rises.

Imagine if the level of η^H is sufficiently low that one can construct an equilibrium with the same structure as in our benchmark setting in which the agent can be induced to exert this level of effort. We can show that in this case the principal always lets the cost of effort rise in the extended model. The intuition for this is that the rise in the cost of effort to the agent serves as an additional form of long run punishment for the agent and therefore provides even better incentives to the agent to exert high effort along the equilibrium path.

More specifically, in the efficient equilibrium of the extended model, the principal chooses the likelihood, intensity, and duration of intervention associated with the level of effort equal to η^H in our benchmark model. Given Proposition 3, this means that the likelihood of intervention is higher, the intensity of intervention is higher, and the duration of intervention is lower compared to the original equilibrium in which the cost of effort does not rise and remains at η^L . Therefore, the level of intensity rises to reinforce the rise in the cost of effort to the agent.

To understand this, note that the first instance of a punishment phase provides the principal with a continuation value of $\underline{J}$ independently of whether the cost of effort to the agent rises or remains the same going forward. Therefore, from an ex-ante perspective, the optimal policy for the principal is to minimize the welfare under a punishment phase for the agent so as to provide the best incentives for the agent to exert effort *along the equilibrium path*. In providing these ex-ante incentives, the principal therefore has two options. One option is to choose $i_t = \tilde{i}$ so as to prevent the cost of effort from rising. The second option is to choose $i^* > \tilde{i}$ and to let the cost of effort rise, where i^* represents the level of intensity which minimizes the agent's welfare from punishment conditional on the cost of effort equal to η^H going forward. It is clear that the principal should choose the second option since, starting from the punishment phase, the agent expects higher levels of intensive force and a higher cost of effort going forward under i^* versus $\tilde{i}$.

Therefore, the long run equilibrium in this extended model features a cooperative and punishment phase which sustain each other as in our benchmark environment, though these are associated with a higher cost of effort to the agent. Moreover, the tradeoff between the intensity and duration of intervention remain and none of our comparative statics change.

As an aside, note that these conclusions change if instead η^H is so high that one cannot construct any equilibrium which sustains high effort by the agent. In this situation, levels of intensity above $\tilde{i}$ cannot be credibly used by the principal since the agent will never exert high effort in the future. Consequently, the optimal punishment for the principal features a cooperative phase and a punishment phase as in our benchmark environment, though the principal sets the level of intensity at $\tilde{i}$ in order to prevent the cost of effort to the agent from rising. Given Proposition 2, this means that there is a higher likelihood of intervention, a lower intensity of intervention, and a longer duration of intervention in comparison to our benchmark environment. Moreover, note that our comparative statics in Proposition 3 must be modified to take into account the fact that the level of intensity does not change with small changes in the environment. Consequently, not only is it the case that the level of intensity does not change, but the duration of intervention actually rises if A declines or if χ rises. This is because, holding the level of intensity constant, these changes enhance the incentives of the principal to punish and hence increase the duration of intervention, and this effect cannot be undone by a rise in intensity as in our benchmark environment.

5.5 Discussion

As discussed in the introduction, there are many applications of our model. A particularly relevant application to current affairs is counterinsurgency policy.⁴³ The majority of modern manuals of counterinsurgency agree that the best way to deal with insurgencies is by obtaining the collaboration of the local leadership.⁴⁴ This means that an analysis of optimal policy under indirect control is particularly relevant. Specifically, the use of costly interventions in this scenario is an important issue in policy discussions. Some experts have defended the use of costly interventions. For example, military strategist Luttwak (2007) writes:

"The simple starting point is that insurgents are not the only ones who

⁴³In this application, the realization of a disturbance corresponds to a successful attack by insurgents. See footnote 11 for how one can model the incentives of the insurgents in our framework.

⁴⁴See Nagl (2002) for a discussion.

can intimidate or terrorize civilians. For instance, whenever insurgents are believed to be present in a village, small town, or distinct city district...the local notables can be compelled to surrender them to the authorities, under the threat of escalating punishments...Occupiers can thus be successful without need of any specialized counterinsurgency methods or tactics if they are willing to out-terrorize the insurgents, so that the fear of reprisals outweighs the desire to help the insurgents or their threats." (p.40-41)

Our model makes three contributions to this policy discussion. First, the model identifies circumstances in which temporary costly interventions—which serve as a form of punishment to the local agent—are optimal. More specifically, it shows that this requires the presence of political economy frictions: double-sided lack of commitment and asymmetric information. It also requires certain additional assumptions. For example, it is necessary that the local agent be more efficient at controlling insurgents relative to the government (Assumption 1) since delegation is otherwise suboptimal. Moreover, it is necessary that the use of excessive force by the government be sufficiently painful to the local agent ($g'(0)$ is sufficiently high) since otherwise temporary costly intervention is suboptimal. Finally, our extensions of Section 5.1 and 5.2 suggest that even if temporary and costly interventions are sometimes optimal, they need only be used *if a sufficient number of disturbances have occurred*. Otherwise, the optimal policy is to provide incentives in the form of rewards, either in the form of payment or in the form of a permanent concessions such as infrastructure investment, political representation, or autonomy.

The second contribution of the model to the policy discussion is that it identifies basic principles that the government should follow while conducting a costly intervention. Importantly, maximal force is inefficient, both because the government must actually use it in equilibrium and also because, if it is too expensive for the government, then it will not be used for sufficiently long. In other words, the government should take into account its own inability to commit to using force. Moreover, the government should use costly intervention as seldomly as possible. What our analysis in Section 4.2 shows is that the optimal contract sets the likelihood of intervention as low as possible so that it is possible for the principal to forgive the agent as often as possible. More specifically, the analysis of Section 4.3 provides precise conditions under which the use of force should be increased or decreased.

The third contribution of the model is that it sheds some light on the role of international pressure against the use of violent interventions (i.e., a rise in A). On the one hand, Proposition 3 states that a government should optimally respond to an increase

in international pressure by reducing intensity i^* , which is the intended consequence of this international pressure. However, on the other hand, Proposition 3 also predicts that an optimally behaving government will also respond with a higher frequency of intervention (higher l^*) and a higher duration of intervention (higher d^*). In sum, international pressure alone cannot remove the need for intervention, and it can have the unintended consequence of making them more frequent and longer. Nonetheless, to the extent that the international community can play a role, the extension in Section 5.2 suggests that one method of actually eradicating equilibrium interventions is to pursue policies which make permanent concessions more desirable than indirect control to the government in question (e.g., setting J^C above $-\pi_a(\eta)\chi/(1-\beta)$ via favors, international concessions, or foreign aid).

6 Conclusion

We have characterized the optimal use of repeated interventions in a model of indirect control. Our explicit closed form solution for the long run dynamics of the efficient sequential equilibrium highlights a fundamental tradeoff between the intensity and duration of interventions. It also allows us to consider the separate effects of a fall in the cost of intensity to the principal, a rise in the cost of disturbances to the principal, and a rise in the cost of effort to the agent.

Our model abstracts from a number of potentially important issues. First, in answering our motivating questions, we have abstracted away from the *static* components of intervention and the means by which a principal directly affects the level of disturbances (i.e., we let π_p be exogenous). Future work should also focus on the static features of optimal intervention and consider how they interact with the dynamic features which we describe. Second, we have ignored the presence of persistent sources of private information. For example, the agent's cost of effort could be unobservable to the principal. Alternatively, the principal may have a private cost of using force. In this latter scenario, a principal with a high cost of force may use more intensive force in order to pretend to have a low cost and to provide better inducements to the agent. We have ignored the presence of persistent hidden information not for realism but for convenience since it maintains the common knowledge of preferences over continuation contracts and simplifies the recursive structure of the efficient sequential equilibria. Understanding the interaction between persistent and temporary hidden information is an important area for future research.

7 Appendix

7.1 Equilibrium Definition

We consider equilibria in which each player conditions his strategy on past public information. Let $h_t^0 = \{z^{1t}, f^{t-1}, z^{2t-1}, i^{t-1}, s^{t-1}\}$, the history of public information at t after the realization of z_t^1 .⁴⁵ Let $h_t^1 = \{h_t^0, f^{t-1}, z^{2t}\}$, the history of public information at t after the realization of z_t^2 . Define a strategy $\sigma = \{\sigma_p, \sigma_a\}$ where $\sigma_p = \{f_t(h_t^0), i_t(h_t^1)\}_{t=0}^\infty$ and $\sigma_a = \{e_t(h_t^1)\}_{t=0}^\infty$ for σ_p and σ_a which are feasible if $f_t(h_t^0) \in \{0, 1\} \forall h_t^0, i_t(h_t^1) \geq 0 \forall h_t^1$, and $e_t(h_t^1) = \{0, \eta\} \forall h_t^1$.

Given σ , define the equilibrium expected continuation values for player j at h_t^0 and h_t^1 , respectively, as $U_j(\sigma|_{h_t^0})$ and $U_j(\sigma|_{h_t^1})$ where $\sigma|_{h_t^0}$ and $\sigma|_{h_t^1}$ correspond to continuation strategies following h_t^0 and h_t^1 , respectively. Let $\Sigma_j|_{h_t^0}$ and $\Sigma_j|_{h_t^1}$ denote the entire set of feasible continuation strategies for j after h_t^0 and h_t^1 , respectively.

Definition 2 σ is a sequential equilibrium if it is feasible and if for $j = p, a$

$$\begin{aligned} U_j(\sigma|_{h_t^0}) &\geq U_j(\sigma'_j|_{h_t^0}, \sigma_{-j}|_{h_t^0}) \quad \forall \sigma'_j|_{h_t^0} \in \Sigma_j|_{h_t^0} \quad \forall h_t^0 \text{ and} \\ U_j(\sigma|_{h_t^1}) &\geq U_j(\sigma'_j|_{h_t^1}, \sigma_{-j}|_{h_t^1}) \quad \forall \sigma'_j|_{h_t^1} \in \Sigma_j|_{h_t^1} \quad \forall h_t^1. \end{aligned}$$

In order to build a sequential equilibrium allocation which is generated by a particular strategy, let $q_t^0 = \{z^{1t}, z^{2t-1}, s^{t-1}\}$ and $q_t^1 = \{z^{1t}, z^{2t}, s^{t-1}\}$, the *exogenous* equilibrium history of public signals and states after the realizations of z_t^1 and z_t^2 , respectively. Define an equilibrium allocation as a function of the exogenous history:

$$\alpha = \{f_t(q_t^0), i_t(q_t^1), e_t(q_t^1)\}_{t=0}^\infty.$$

Let $\mathcal{F}$ denote the set of feasible allocations α with continuation allocations from t onward which are measurable with respect to public information generated up to t . Let $U_j(\alpha|_{q_t^0})$ and $U_j(\alpha|_{q_t^1})$ correspond to the equilibrium continuation value to player j following the realization of q_t^0 and q_t^1 , respectively. The following lemma provides necessary and sufficient conditions for α to be generated by sequential equilibrium strategies.

⁴⁵Without loss of generality, we let $i_t = 0$ if $f_t = 0$ and $e_t = 0$ if $f_t = 1$.

Lemma 1 $\alpha \in \mathcal{F}$ is a sequential equilibrium allocation if and only if

$$U_j(\alpha|_{q_t^0}) \geq \underline{U}_j \text{ for } j = p, a \ \forall q_t^0, \quad (14)$$

$$U_p(\alpha|_{q_t^1}) \geq -\pi_p \chi + \beta \underline{U}_p \ \forall q_t^1 \text{ s.t. } f_t(q_t^0) = 1, \text{ and} \quad (15)$$

$$U_a(\alpha|_{q_t^1}) \geq \max \left\{ \begin{array}{l} -\eta + \beta \left(\begin{array}{l} (1 - \pi_a(\eta)) \left\{ U_j(\alpha|_{q_{t+1}^0}) | q_t^1, s_t = 0 \right\} \\ + \pi_a(\eta) E \left\{ U_j(\alpha|_{q_{t+1}^0}) | q_t^1, s_t = 1 \right\} \end{array} \right) \\ \beta \left(\begin{array}{l} (1 - \pi_a(0)) \left\{ U_j(\alpha|_{q_{t+1}^0}) | q_t^1, s_t = 0 \right\} \\ + \pi_a(0) E \left\{ U_j(\alpha|_{q_{t+1}^0}) | q_t^1, s_t = 1 \right\} \end{array} \right) \end{array} \right\} \ \forall q_t^1 \text{ s.t. } f_t(q_t^0) = 0 \quad (16)$$

for $\underline{U}_p = -\pi_p \chi / (1 - \beta)$ and some $\underline{U}_a \leq w_a / (1 - \beta)$.

Proof. Step 1. The necessity of (14) for $j = p$ follows from the fact that the principal can choose $f'_k(q_k^0) = 1 \ \forall k \geq t$ and $\forall q_k^0$ and $i'_k(q_k^1) = 0 \ \forall k \geq t$ and $\forall q_k^1$, and this delivers continuation value $\underline{U}_p$. The necessity of (14) for $j = a$ follows from the fact that the agent can choose $e'_k(q_k^1) = 0 \ \forall k \geq t$ and $\forall q_k^1$, and this delivers a continuation value of at least $\underline{U}_a$. **Step 2.** The necessity of (15) follows from the fact that conditional on $f_t(q_t^0) = 1$, the principal can choose $f'_k(q_k^0) = 1 \ \forall k > t$ and $\forall q_k^0$ and $i'_k(q_k^1) = 0 \ \forall k \geq t$ and $\forall q_k^1$, and this delivers continuation value $-\pi_p \chi + \beta \underline{U}_p$. The necessity of (16) follows from the fact that conditional on $f'_t(q_t^0) = 0$, the agent can unobservably choose $e'_t(q_t^1) \neq e_t(q_t^1)$ and follow the equilibrium strategy $\forall k > t$ and $\forall q_k^1$. **Step 3.** For sufficiency, consider a feasible allocation which satisfies (14) – (16) and construct the following off-equilibrium strategy. Any observable deviation by the principal results in a reversion to the repeated static Nash equilibrium. We only consider single period deviations since $\beta < 1$ and since continuation values are bounded. If $f'_t(q_t^0) = 1$, then a deviation by the principal to $f'_t(q_t^0) = 0$ is weakly dominated by (14) and Assumption 2. Moreover, a deviation by the principal to $i'(q_t^1) \neq i(q_t^1)$ is weakly dominated by (15). If $f_t(q_t^0) = 0$, then a deviation by the principal to $f'_t(q_t^0) = 1$ is weakly dominated by (14). If $f_t(q_t^0) = 0$, then a deviation by the agent to $e'_t(q_t^1) \neq e_t(q_t^1)$ is weakly dominated by (16). ■

7.2 Implications of Assumption 3

The value of $\hat{\beta}$ satisfies

$$\widehat{\beta} = \max \left\{ \frac{\eta}{\eta(1 - \pi_a(0)) - w_a(\pi_a(0) - \pi_a(\eta))}, \frac{1}{A} \frac{g'(i^*) i^*}{-\frac{\pi_a(0)}{\pi_a(0) - \pi_a(\eta)} \eta - w_a + g(i^*)} \right\}$$

for i^* which satisfies (9). Given the functions $l(i)$ and $d(i)$ defined in Section 4.2, the first part of this assumption implies that $l(0) < 1$ so that an equilibrium in which high effort is sustained by the threat of the repeated static Nash equilibrium exists. Since $l(i)$ is declining in i for $i < i^*$ by Proposition 2, this assumption guarantees that $l(i^*) < 1$. The second part of this assumption implies that $d(i^*) > 0$. These features guarantee that the set of values $U \in [\underline{U}, \overline{U}]$ are self-generating so that the long run equilibrium can be explicitly characterized.

7.3 Proofs of Additional Lemmas

In this section we prove several important lemmas which are required for proving our propositions. Let Γ represent the set of sequential equilibrium continuation values and let $U^{\max}$ the highest continuation value to the agent in this set.

Lemma 2 (i) Γ is convex and compact, (ii) $J(\underline{U}) = J(U^{\max}) = \underline{J}$, and (iii) $J(U)$ is weakly concave.

Proof. Step 1. The weak concavity of the program and the convexity of the constraint set in (1) – (7) guarantees that Γ is convex. **Step 2.** If we set an arbitrarily high upper bound for i in (1) – (7), then the compactness of the constraint set together with the fact that $\beta < 1$ guarantees that Γ is closed and bounded by the Dominated Convergence Theorem. **Step 3.** By (3), $J(\underline{U}) \geq \underline{J}$ and $J(U^{\max}) \geq \underline{J}$. **Step 4.** By Assumptions 1 and 2 and equations (4) and (6), it must be that $f_z^*(\underline{U}) = 1 \forall z$ since otherwise an increase in f_z for some z must satisfy (3) – (7) and strictly reduces the welfare of the agent. If $J(\underline{U}) > \underline{J}$, then an increase i_z must satisfy (3) – (7) and strictly reduces the welfare of the agent. Therefore $J(\underline{U}) = \underline{J}$. **Step 5.** By Assumption 1 and equations (4) and (6), $f_z^*(U^{\max}) = 0 \forall z$ since otherwise a decrease in f_z for some z must satisfy (3) – (7) and strictly increase the welfare of the agent. If $J(U^{\max}) > \underline{J}$, then a decrease in e_z or an increase in U_z^H strictly increases the welfare of the principal while satisfying (3) – (7), and if this were not feasible then $U^{\max} = 0$, which violates (3) since it implies $J(U^{\max}) = -\pi_a(0)\chi$. Therefore, $J(U^{\max}) = \underline{J}$. **Step 6.** The weak concavity of $J(\cdot)$ follows directly from the first and second parts of the lemma. ■

Lemma 3 $\exists i^*$ s.t. the solution to (1) – (7) cannot admit $i_z^*(U) \neq i^*$ for any z given $f_z^*(U) = 1$.

Proof. Step 1. Define $i^* = Ei_z^*(\underline{U})$. By Step 4 of the proof of Lemma 2, $f_z^*(\underline{U}) = 1 \forall z$. It must be that $i_z^*(\underline{U}) = i^* \forall z$ since otherwise a perturbation which sets $i_z^*(\underline{U}) = Ei_z^*(\underline{U}) \forall z$ continues to satisfy (3) – (7) and strictly reduces the welfare of the agent by the concavity of $g(i)$ and $J(U)$. **Step 2.** Let $\hat{J}(U|\hat{i})$ correspond to the maximizer of (1) – (7) subject to the additional constraints that $f_z = 1$ and $i_z = \hat{i} \forall z$ for some $\hat{i}$. Note that for any two value U' and U'' where $(w_a - g(\hat{i})) / (1 - \beta) \leq U' < U''$, it must be that

$$\frac{\hat{J}(U''|\hat{i}) - \hat{J}(U'|\hat{i})}{U'' - U'} = \frac{J\left(\frac{U'' - w_a + g(\hat{i})}{\beta}\right) - J\left(\frac{U' - w_a + g(\hat{i})}{\beta}\right)}{\frac{U'' - U'}{\beta}} \quad (17)$$

$$\leq \frac{J(U'') - J(U')}{U'' - U'}, \quad (18)$$

where we have appealed to the concavity of $J(\cdot)$. **Step 3.** Imagine if $\exists \hat{i} \neq i^*$ s.t. $\hat{J}(U|\hat{i}) = J(U)$ for some U . Let $\hat{U}(\hat{i}) \geq (w_a - g(\hat{i})) / (1 - \beta)$ denote the value which solves $\hat{J}(\hat{U}(\hat{i})|\hat{i}) = \underline{J}$ for such $\hat{i}$, which must exist by the concavity of $\hat{J}(\cdot)$ since $\hat{J}(U|\hat{i}) \geq \underline{J}$ for some U . By step 1 and Assumption 3, $J(\hat{U}(\hat{i})) > \hat{J}(\hat{U}(\hat{i})|\hat{i})$, so that by (18) $\hat{J}(U|\hat{i}) < J(U) \forall U \geq \hat{U}(\hat{i})$. Therefore, $\hat{J}(U|\hat{i}) < J(U) \forall U$ and $\forall \hat{i} \neq i^*$. **Step 4.** By step 3, $i_z^*(U) = i^*$ if $f_z^*(U) = 1 \forall z$. ■

Lemma 4 $\exists \tilde{U} \in (\underline{U}, U^{\max})$ and some $m > 0$ s.t.

$$f_z^*(U) = 0 \forall z \text{ and } \forall U \geq \tilde{U} \text{ and}$$

$$J(U) \begin{cases} = \underline{J} + m(U - \underline{U}) & \text{if } U \leq \tilde{U}. \\ < \underline{J} + m(U - \underline{U}) & \text{if } U > \tilde{U} \end{cases}$$

Proof. Step 1. Consider two continuation values $U' < U''$ s.t. $Ef_z^*(U') > 0$ and $Ef_z^*(U'') > 0$. It follows given Lemma 3 that

$$J(U) = J(U') + m(U - U') \quad \forall U \in [U', U''] \quad (19)$$

$$\text{where } m = \frac{J(U'') - J(U')}{U'' - U'}.$$

To see why, let $U^{W*}(U)$ correspond to the expected continuation value to the agent conditional on $f_z = 1$ and let $U^{P*}(U)$ correspond to the expected continuation value to the agent conditional on $f_z = 0$. Optimality and the concavity of $J(\cdot)$ thus require

$$J(U) = J(U^{W*}(U)) Ef_z^*(U) + J(U^{P*}(U)) (1 - Ef_z^*(U)). \quad (20)$$

By (20) and the concavity of $J(\cdot)$, it follows that $U^{W*}(U)$ and $U^{P*}(U)$ are on the same line segment in $J(\cdot)$ for a given U . By the concavity of $J(\cdot)$, one can choose $\forall z, U_z^{F*}(U^{W*}(U)) = \frac{U^{W*}(U) - w_a + g(i^*)}{\beta} \geq U^{W*}(U)$ which is weak if $i^* > 0$, so that

$$\frac{J(U^{W*}(U'')) - J(U^{W*}(U'))}{U^{W*}(U'') - U^{W*}(U')} = \frac{\left(J\left(\frac{U^{W*}(U'') - w_a + g(i^*)}{\beta}\right) - J\left(\frac{U^{W*}(U') - w_a + g(i^*)}{\beta}\right) \right)}{\frac{U^{W*}(U'') - U^{W*}(U')}{\beta}}. \quad (21)$$

By the concavity of $J(\cdot)$, this implies $U^{W*}(U'')$ and $U^{W*}(U')$ are on the same line segment. Therefore, (19) applies. **Step 2.** Since $Ef_z^*(\underline{U}) = 1$ by step 3 of the proof of Lemma 2, it follows from step 1 that (19) applies for $U' = \underline{U}$ and some $U'' = \tilde{U} \geq \underline{U}$. It follows that $f_z^*(U) = 0 \forall z$ and $\forall U \geq \tilde{U}$ if $\tilde{U} > \underline{U}$ and $f_z^*(U) = 0 \forall z$ and $\forall U > \tilde{U}$ if $\tilde{U} = \underline{U}$. **Step 3.** If $\tilde{U} = \underline{U}$, then $Ef_z^*(\underline{U}) = 0 \forall U > \underline{U}$, but this is not possible since (2) and (6) imply that $EU_z^{L*}(U) < U$ and cannot be arbitrarily close to U . Therefore $m > 0$. **Step 4.** It cannot be that $\tilde{U} = U^{\max}$ since this violates part 2 of Lemma 2. ■

Lemma 5 $\tilde{U} = \bar{U}$.

Proof. Step 1. $e_z^*(U) = \eta$ if $f_z^*(U) = 0$ and $U \in [\underline{U}, \tilde{U}]$. Suppose this is not the case and consider a solution for which $e_z^*(U) = 0$ and $f_z^*(U) = 0$. Because the constraint set is convex, one can perturb this solution without changing welfare so that (6) binds and $U_z^{L*}(U) = U_z^{H*}(U)$. However, this implies that $U_z^{L*}(U) < -e_z^*(U) + \beta U_z^{L*}(U)$. Because optimality given the concavity of $J(\cdot)$ requires $-e_z^*(U) + \beta U_z^{L*}(U) \in [\underline{U}, \tilde{U}]$, this means given Lemma 4 that $-\pi_a(0)\chi + \beta J(U_z^{L*}(U)) < \underline{J}$, which violates (3). **Step 2.** Suppose $\bar{U} < \tilde{U}$. By Assumption 3, there exists a solution to (1) – (7) s.t. $f_z^*(U) = 0$ and $e_z^*(U) = \eta \forall z$ and $\forall U \in [\bar{U}, \tilde{U}]$. Moreover, given the concavity of the program and convexity of the constraint set in (1) – (7) such a solution can feature $U_z^{H*}(U) = U^{H*}(U)$ and $U_z^{L*}(U) = U^{L*}(U) \forall z$. This implies that

$$m = \frac{J(\tilde{U}) - J(\bar{U})}{\tilde{U} - \bar{U}} = (1 - \pi_a(\eta)) \frac{J(U^H(\tilde{U})) - J(U^H(\bar{U}))}{U^H(\tilde{U}) - U^H(\bar{U})} + \pi_a(\eta) m,$$

but since $U^H(\tilde{U}) > \tilde{U}$, this violates Lemma 4. **Step 3.** Suppose $\bar{U} > \tilde{U}$ so that by Lemma 4, $J(\tilde{U} + \epsilon) < \underline{J} + m(U - \underline{U})$ for $\epsilon > 0$ arbitrarily small. Consider a perturbation which sets $e_z^*(\tilde{U} + \epsilon) = \eta$ and lets (6) bind so that $U_z^{L*}(\tilde{U} + \epsilon) < U_z^{H*}(\tilde{U} + \epsilon) < \tilde{U} \forall z$. This perturbation yields a payoff to the principal equal to $\underline{J} + m(\tilde{U} + \epsilon - \underline{U})$, violating the definition of $\tilde{U}$ in Lemma 4. ■

7.4 Proof of Proposition 1

Step 1. We begin by characterizing the solution for $U \in [\underline{U}, \bar{U}]$ to prove the second part of the proposition and having done this we prove the first part of the proposition. By steps 4 and 5 of the proof of Lemma 2 and by Lemma 4, the solution which satisfies the *Bang-Bang* property is characterized by a probability $E f_z^*(U) = (\bar{U} - U) / (\bar{U} - \underline{U})$, where

$$\begin{aligned}\underline{U} &= E \{ w_a - g(i_z^*(U)) + \beta U_z^{F*}(U) | f_z^*(U) = 1 \} \\ \bar{U} &= E \{ -e_z^*(U) + \beta ((1 - \pi_a(e_z^*(U))) U_z^{H*}(U) + \pi_a(e_z^*(U)) U_z^{L*}(U)) | f_z^*(U) = 0 \},\end{aligned}$$

and the analogous expected continuation values for the principal are $J(\bar{U})$ and $J(\underline{U}) = \underline{J}$, respectively. Therefore, one only needs to characterize $i_z^*(\underline{U})$, $e_z^*(\bar{U})$, $U_z^{F*}(\underline{U})$, $U_z^{H*}(\bar{U})$, and $U_z^{L*}(\bar{U})$ to achieve a full description of equilibrium actions. **Step 2.** By Lemma 3 $i_z^*(\underline{U}) = i^* \forall z$. By step 2 of the proof of Lemma 5, $e_z^*(\bar{U}) = \eta \forall z$. **Step 3.** By Lemmas 4 and 5 $U_z^{H*}(\bar{U}) = \bar{U}$ and $U_z^{L*}(U) = \bar{U} - \eta / (\beta(\pi_a(0) - \pi_a(\eta))) < \bar{U}$ since otherwise (6) does not bind and a perturbation which reduces U_z^H and raises U_z^L strictly raises welfare. **Step 4.** The fact that $U_z^{F*}(\underline{U}) = (\underline{U} - w_a + g(i^*)) / \beta \forall z$ is implied by (2) and the fact that (5) binds since otherwise the principal is receiving a continuation value above $\underline{J}$. **Step 5.** We are left to characterize i^* and $\underline{U}$. Note that the equilibrium can be represented by a system of 4 equations: (10) – (12) and

$$\underline{J} \leq -\pi_p \chi - A i^* + \beta ((1 - d^*) J(\bar{U}) + d^* \underline{J}). \quad (22)$$

(22) is an equality if $d^* \geq 0$ which occurs if $(\underline{U} - w_a + g(i^*)) / \beta \leq \bar{U}$, where we have taken Lemma 4 into account. (10) – (12) and (22) represent a system of 4 equations and 5 unknowns: $J(\bar{U})$, $\underline{U}$, l^* , i^* , and d^* , where the fifth unknown is pinned down by the fact that these variables are chosen to maximize $J(\bar{U})$. Note that given steps 1-4, $d^* < 1$ and

$l^* \in (0, 1)$ so that by algebraic substitution, it is the case that

$$J(\bar{U})(1 - \beta) \leq \frac{-\frac{\pi_a(0)}{\pi_a(0) - \pi_a(\eta)}\eta - w_a + g(i^*)}{-\eta - w_a + g(i^*)} ((\pi_p - \pi_a(\eta))\chi + Ai^*) - (\pi_p\chi + Ai^*), \quad (23)$$

which is an equality if and only if (22) is an equality. **Step 6.** Note that i^* which satisfies (9) maximizes the right hand side of (23). Moreover, by Assumption 3, it is the case in the optimum that (22) binds since the implied value of d^* exceeds 0 so that $U_z^{F*}(\underline{U}) = (\underline{U} - w_a + g(i^*)) / \beta \leq \bar{U}$. Substitution into (11) yields $\underline{U}$ which completes the proof of the second part. **Step 7.** Lemmas 4 and 5 imply that if $\Pr\{U_t \geq \bar{U} \forall t\} > 0$, then $\Pr\{f_t = 0 \forall t\} > 0$. However, (2) and (6) imply that $\Pr\{U_{t+1} < U_t - \epsilon | f_t = 0\} > 0 \forall t$ for some $\epsilon > 0$, which means that $\Pr\{U_t \geq \bar{U} \forall t\} = 0$. **Step 8.** $\Pr\{U_{t+1} \leq \bar{U} | U_t \leq \bar{U}\} = 1$ by steps 3 and 6, so that by step 7, $\lim_{t \rightarrow \infty} \Pr\{U_t \leq \bar{U}\} = 1 \forall U_0$. **Q.E.D.**

7.5 Proof of Proposition 2

Step 1. An equilibrium with the given structure satisfies (10) – (13) and entails functions $l(i)$ and $d(i)$ defined by:

$$\begin{aligned} 1 - \pi_a(\eta)l(i) &= \frac{\gamma_a + \beta\gamma_p - 1}{\beta(\gamma_a + \gamma_p - 1)} \\ d(i) &= \frac{\gamma_p + \beta\gamma_a - 1}{\beta(\gamma_a + \gamma_p - 1)} \end{aligned}$$

for

$$\gamma_a = \frac{-\frac{\pi_a(0)}{\pi_a(0) - \pi_a(\eta)}\eta - w_a + g(i)}{-\eta - w_a + g(i)} \quad (24)$$

$$\gamma_p = \frac{(\pi_p - \pi_a(\eta))\chi}{(\pi_p - \pi_a(\eta))\chi + Ai} \quad (25)$$

where Assumption 3 and the fact that $\beta < 1$ implies $\gamma_a \in [0, 1]$, $\gamma_p \in [0, 1]$, and $\gamma_a + \gamma_p - 1 > 0$ for all $i \leq \bar{i}$, where $\bar{i} > i^*$. **Step 2.** By some algebra, $l'(i)$ has the same sign as $-\gamma_p \partial \gamma_a / \partial i - (1 - \gamma_a) \partial \gamma_p / \partial i$ which equals

$$-\gamma_p(1 - \gamma_a) \left(\frac{g'(i)}{-\eta - w_a + g(i)} - \frac{A}{(\pi_p - \pi_a(\eta))\chi + Ai} \right). \quad (26)$$

Since $g(\cdot)$ is concave, it follows that (26) is negative if $i < i^*$ and positive if $i > i^*$. **Step 3.** By some algebra, $d'(i)$ has the same sign as $(1 - \gamma_p) \partial \gamma_a / \partial i + \gamma_a \partial \gamma_p / \partial i$ which equals

$$-\frac{A}{[-\eta - (\omega_a - g(i))][(\pi_a(\eta) - \pi_p)\chi - Ai]} \times \left[ig'(i)(1 - \gamma_a) - \left[\frac{-\eta \pi_a(0)}{\pi_a(0) - \pi_a(\eta)} - (\omega_a - g(i)) \right] \gamma_p \right] \quad (27)$$

The element outside the square brackets is always positive. Consider $i < i^*$, where the concavity of $g(\cdot)$ guarantees that

$$-g'(i)[(\pi_a(\eta) - \pi_p)\chi - Ai] - A(-\eta - (\omega_a - g(i))) > 0. \quad (28)$$

By some algebra, one can show that given (28), the element inside the square brackets in (27) is decreasing in i for $i < i^*$. Since $d(0) = 1$, it follows that $d'(0) \leq 0$, so this fact implies that $d'(i) < 0$ for $i < i^*$. Consider $i \geq i^*$. By rearranging terms, $(1 - \gamma_p) \partial \gamma_a / \partial i + \gamma_a \partial \gamma_p / \partial i$ can also be expressed as

$$(1 - \gamma_p - \gamma_a) \frac{g'(i)}{-\eta - (\omega_a - g(i))} + \gamma_a \gamma_p \left[\frac{g'(i)}{-\eta - (\omega_a - g(i))} + \frac{A}{(\pi_a(\eta) - \pi_p)\chi - Ai} \right],$$

which is negative for $i \geq i^*$ since the left hand side of (28) is weakly negative in this case.

Q.E.D.

7.6 Proof of Proposition 3

Step 1. Implicit differentiation of (9) taking into account the concavity of $g(\cdot)$ yields the comparative statics with respect to i^* . **Step 2.** Given γ_a and γ_p defined in the proof of Proposition 2, it is the case that if a particular parameter $x = \{A, \chi, \eta\}$ changes, the effect on l^* has the same sign as

$$-\gamma_p \frac{\partial \gamma_a}{\partial x} - (1 - \gamma_a) \frac{\partial \gamma_p}{\partial x}, \quad (29)$$

where we have used the fact that $-\gamma_p \partial \gamma_a / \partial i - (1 - \gamma_a) \partial \gamma_p / \partial i = 0$ at i^* . The effect on d^* has the same sign as

$$\left((1 - \gamma_p) \frac{\partial \gamma_a}{\partial x} + \gamma_a \frac{\partial \gamma_p}{\partial x} \right) + \left((1 - \gamma_p) \frac{\partial \gamma_a}{\partial i} + \gamma_a \frac{\partial \gamma_p}{\partial i} \right) \frac{\partial i}{\partial x}. \quad (30)$$

Step 3. Given (29), the comparative statics with respect to l^* are implied by the fact that $\partial\gamma_a/\partial A = \partial\gamma_a/\partial\chi = 0$, $\partial\gamma_a/\partial\eta < 0$, $\partial\gamma_p/\partial A < 0$, $\partial\gamma_p/\partial\chi > 0$, and $\partial\gamma_p/\partial\eta = 0$. **Step 4.** Given (30), the effect of an increase in η on d^* is implied by the fact that $\partial\gamma_a/\partial\eta < 0$, $\partial\gamma_p/\partial\eta = 0$, $\partial i/\partial\eta > 0$, and $(1 - \gamma_p) \partial\gamma_a/\partial i + \gamma_a \partial\gamma_p/\partial i < 0$ from Proposition 2. Letting $x = A$, substitution into (30) taking Assumption 4 into account together with $\partial i/\partial A = \gamma_p g'(i) / Ag''(i)$ yields:

$$\left(\frac{\gamma_a \gamma_p - (1 - \gamma_p)(1 - \gamma_a)}{\gamma_p} \right) \left(\frac{\partial\gamma_p}{\partial i} \right) \left(-\gamma_p \frac{i}{A} \frac{1}{(1 - \theta)} \right) + \gamma_a \frac{\partial\gamma_p}{\partial i} \frac{i}{A}$$

which has the same sign as $\theta\gamma_a + \gamma_p - 1$, which is unambiguously positive given the definition of i^* . Analogous arguments imply the comparative static with respect to χ . **Q.E.D.**

8 Bibliography

Abreu, Dilip (1988) "On the Theory of Infinitely Repeated Games with Discounting," *Econometrica*, 56, 383-397.

Abreu, Dilip, David Pearce, and Ennio Stacchetti (1990) "Towards a Theory of Discounted Repeated Games with Imperfect Monitoring," *Econometrica*, 58, 1041-1063.

Acemoglu, Daron, Mikhail Golosov, and Aleh Tsyvinski (2008) "Political Economy of Mechanisms," *Econometrica*, 76, 619-641.

Acemoglu, Daron and James A. Robinson (2006), *Economic Origins of Dictatorship and Democracy*, Cambridge University Press; New York.

Acemoglu, Daron and Alexander Wolitzky (2009) "The Economics of Labor Coercion," Working Paper.

Albuquerque, Rui and Hugo Hopenhayn (2002) "Optimal Lending Contracts and Firm Dynamics," *Review of Economic Studies*, 71, 285-315.

Anderlini, Luca, Dino Gerardi, and Roger Lagunoff (2009) "Social Memory, Evidence, and Conflict," forthcoming in *Review of Economic Dynamics*.

Atkeson, Andrew and Robert Lucas (1992) "On Efficient Distribution with Private Information," *Review of Economic Studies*, 59, 427-453.

Baliga, Sandeep and Tomas Sjöström (2004) "Arms Races and Negotiations," *Review of Economic Studies*, 71, 351-369.

Chassang, Sylvain and Gerard Padró i Miquel (2009) "Economic Shocks and Civil War," Working Paper.

Chwe, Michael Suk-Young (1990) "Why Were Workers Whipped? Pain in a Principal-Agent Model," *Economic Journal*, 100, 1109-1121

Dal Bó, Ernesto, Pedro Dal Bó, and Rafael Di Tella (2006) "Plata o Plomo?: Bribe and Punishment in a Theory of Political Influence," *American Political Science Review*, 100, 41-53.

Dal Bó, Ernesto and Rafael Di Tella (2003) "Capture by Threat," *Journal of Political Economy*, 111, 1123-1154, .

Debs, Alexandre (2008) "Political Strength and Economic Efficiency in a Multi-Agent State," Working Paper.

Egorov, Georgy and Konstantin Sonin (2009) "Dictatorships and their Viziers: Endogenizing the Loyalty-Competence Tradeoff," forthcoming in *Journal of the European Economic Association*.

Esteban, Joan and Debraj Ray (2008) "On the Saliency of Ethnic Conflict," forthcoming, *American Economic Review*.

- Fearon, James D. (1995)** "Rationalist Explanations for War," *International Organization*, 49, 379-414.
- Ferejohn, John (1986)** "Incumbent Performance and Electoral Control," *Public Choice*, 50, 5-25.
- Fong, Yuk-fai and Jin Li (2009)** "Relational Contracts, Limited Liability, and Employment Dynamics," Working Paper.
- Fudenberg, Drew, David Levine, and Eric Maskin (1994)** "The Folk Theorem with Imperfect Public Information," *Econometrica*, 62, 997-1039.
- Golosov, Mikhail, Narayana Kocherlakota, and Aleh Tsyvinski (2003)** "Optimal Indirect and Capital Taxation," *Review of Economic Studies*, 70, 569-587.
- Green, Edward J. and Robert H. Porter (1984)** "Noncooperative Collusion Under Imperfect Price Information," *Econometrica*, 52, 87-100.
- Herbst, Jeffrey (2000)** *States and Power in Africa*, Princeton University Press, Princeton.
- Jackson, Matthew O. and Massimo Morelli (2008)** "Strategic Militarization, Deterrence, and War," Working Paper.
- Luttwak, Edward N. (1976)** *The Grand Strategy of the Roman Empire*, Johns Hopkins University Press, Baltimore and London.
- Luttwak, Edward N. (2007)** "Dead End: Counterinsurgency Warfare as Military Malpractice," *Harper's Magazine*, February, 33-42.
- Myerson, Roger (2008)** "Leadership, Trust, and Power: Dynamic Moral Hazard in High Office," Working Paper.
- Nagl, John A. (2002)** *Learning to Eat Soup with a Knife: Counterinsurgency Lessons from Malaya and Vietnam*, Chigaco University Press, Chicago.
- Phelan, Christopher (1995)** "Repeated Moral Hazard and One-Sided Commitment," *Journal of Economic Theory*, 66, 488-506.
- Polinsky, A. Mitchell and Steven Shavell (1979)** "The Optimal Tradeoff between the Probability and Magnitude of Fines," *American Economic Review*, 69, 880-891.
- Polinsky, A. Mitchell and Steven Shavell (1984)** "The Optimal Use of Fines and Imprisonment," *Journal of Public Economics*, 24, 89-99.
- Powell, Robert (1999)** *In the Shadow of Power: States and Strategies in International Politics*, Princeton University Press; Princeton.
- Reno, William (1998)** *Warlords Politics and African States*, Lynne Rienner Publishers; London.
- Syme, Ronald (1933)** "Some Notes on the Legions under Augustus," *Journal of Roman Studies*, 23: 14-33.

Schwarz, Michael and Konstantin Sonin (2004) "A Theory of Brinksmanship, Conflicts, and Commitments," *Journal of Law, Economics, and Organization*, forthcoming.

Thomas, Jonathan and Tim Worrall (1990) "Income Fluctuation and Asymmetric Information: An Example of a Repeated Principal-Agent Problem," *Journal of Economic Theory* 51, 367-390.

Yared, Pierre (2009) "A Dynamic Theory of War and Peace", Working Paper.