

The Political Economy of Indirect Control*

Gerard Padró i Miquel[†]

Pierre Yared[‡]

August 2, 2011

Abstract

This paper characterizes optimal policy when 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 repeated moral hazard model with two additional features. First, the principal is allowed to exert direct control by intervening with an endogenously determined intensity of force which is costly to both players. Second, the principal suffers from limited commitment. Using recursive methods, we derive a fully analytical characterization of the intensity, likelihood, and duration of intervention. The first main insight from our model is that repeated and costly equilibrium interventions are a feature of optimal policy. This is because they are the most efficient credible means for the principal of providing incentives for the agent. 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 intensity, likelihood, and duration of intervention. We discuss these results in the context of some historical episodes.

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

JEL Classification: D02, D82, H1

*We would like to thank Daron Acemoglu, Effi Benmelech, Carmen Beviá, Claudine Desrieux, Dennis Gromb, Mike Golosov, Marina Halac, Elhanan Helpman, Johannes Horner, Narayana Kocherlakota, Gilat Levy, Jin Li, Wolfgang Pesendorfer, Robert Powell, Nancy Qian, Ronny Razin, Dmitriy Sergeyev, Ali Shourideh, Kjetil Storesletten, Aleh Tsyvinski, Joachim Voth, five anonymous referees, and seminar participants at INSEAD, Kellogg MEDS, Minneapolis Fed, Paris School of Economics, Princeton Political Economy Conference, and the Social Determinants of Conflict Conference for comments. Gerard Padró i Miquel gratefully acknowledges financial support from ESRC under grant RES-061-250170. All remaining mistakes are our own.

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

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

1 Introduction

In exerting their authority in weakly institutionalized environments, governments often use indirect control: Certain political responsibilities are left to local authorities or warlords who have an advantage in fulfilling them. These tasks range from the prevention of riots, protests, and crime, to 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 over a vast territory with great economy of force. These clients were kept in line by a combination of subsidies and favors and by the threat of military intervention.¹ Beyond Roman times, this strategy of indirect control through violent interventions 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 Many of these interventions are temporary, repeated, and often deemed excessively destructive.

In this paper, we ask the following question: How should a government use rewards and military interventions to align the incentives of the local authority with its own? More specifically, when should interventions be used? How long and how intense should they be? In answering these questions, it is important to take into account that the interaction between a government and a local authority is inherently dynamic, and that there are three key frictions to consider.

First, the local authority cannot commit to fulfilling its delegated task. Second, the local authority's actions, which often occur through informal channels, are imperfectly observed by the government. Third, the government cannot commit to providing rewards or using interventions. While the first two constraints point to a classic moral hazard problem, the optimal policy in this context must take into account how the third constraint interacts with the first two. Therefore, a suitably modified repeated principal-agent model (in which the government is the principal) is the natural framework to provide guidance on the implications of these frictions for optimal policy.

In this paper, we develop such a model. The principal delegates the control of disturbances—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 repeated moral hazard model with two additional features that are crucial in our application.³ First, the principal is allowed to intervene with an endogenously determined intensity of force that is costly to both players. Second, the principal suffers from limited commitment. We focus on characterizing the optimal intensity, likelihood, and duration of such interventions. Using the recursive methods of Abreu, Pearce,

¹See Syme (1933), Luttwak (1976), and our discussion in Section 7 for more details.

²This is particularly the case for governments that have tenuous control over parts of their territory, for instance, in Pakistan's Federally Administered Tribal Areas and in outlying areas in many African countries. On this point, see Herbst (2000) and Reno (1998). Recent violent 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.

³See Debs (2009), Egorov and Sonin (2009), Guriev (2004), and Myerson (2008) for applications of a principal-agent model to delegation problems in weakly institutionalized environments such as dictatorships.

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 equilibrium interventions are a feature of optimal policy and occur along the equilibrium path.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 important factors on the optimal intensity, likelihood, and duration of intervention.

More specifically, we construct a repeated game between a principal and an agent. In every period, the principal has two options. On the one hand, the principal can withhold force and allow the agent to exert *unobservable* effort in controlling disturbances. In this situation, if a large disturbance occurs, the principal cannot determine if it is due to the agent's negligence or due to bad luck. On the other hand, the principal can directly intervene to control disturbances himself, and in doing so he chooses the intensity of force, where higher intensity hurts both the principal and the agent. Both the principal and the agent suffer from limited commitment. Because the agent cannot commit to high effort once the threat of intervention has subsided, the Nash equilibrium of the stage game is direct intervention by the principal with minimal intensity (i.e., direct control). Dynamic incentives, however, can generate better outcomes in which the agent exerts effort. We consider the efficient sequential equilibrium of this game in which reputation sustains equilibrium actions, and we fully characterize in closed form the dynamics of interventions.

Our first result is that repeated and costly interventions are a feature of optimal policy since they are the most efficient credible means for the principal of providing incentives for the agent. Interventions must occasionally be used following large disturbances as a costly punishment to induce the agent to exert high effort along the equilibrium path. Moreover, these interventions must be temporary because of limited commitment on the side of the principal. We show that if the principal could commit to a long run contract, then optimal interventions would last forever in order to provide the best ex-ante incentives for the agent. The principal's inability to commit implies that any costly intervention must be followed by periods of cooperation in which the exertion of effort by the agent rewards the principal for having intervened.

More specifically, we show that once the first intervention takes place, the principal and the agent engage in 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 large

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)

disturbance can trigger a transition to the punishment phase. In the punishment phase which follows a large disturbance, 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, and the principal prefers to maintain high effort by the agent in the future. Since the punishment phase is costly to both players, the optimal contract minimizes the likelihood of transitioning to this phase. Hence, the principal must provide the strongest incentives for the agent to exert high effort during the cooperative phase. This is achieved by the principal promising the harshest credible punishment to the agent in case disturbances are very large. Thus, the worst credible punishment to the agent sustains the highest welfare for both the principal and the agent during phases of cooperation by minimizing the likelihood of punishment.

Our second result follows from our explicit characterization of the worst credible punishment to the agent. Recall that the principal cannot commit to future actions. This inability to commit 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 of force if cooperation is expected to resume in the future, and higher intensity is only credible if cooperation resumes sooner. This trade-off between intensity and duration generates a non-monotonic relationship between the intensity chosen by the principal and the agent's overall payoff during punishment phases. At low levels of intensity, the agent's payoff naturally becomes worse as intensity rises. However, at higher levels of intensity, the marginal instantaneous cost to the agent from higher intensity is counteracted by the shorter duration of punishment. Hence, the overall punishment phase is less harsh as intensity further increases. The principal takes into account these opposing forces to determine the worst credible punishment to the agent.

Our final result concerns the effect of three important factors on the optimal intensity, likelihood, and duration of intervention. First, we consider the effect of a rise in the cost of effort to the agent. Second, we consider the effect of a decline in the cost of intensity to the principal. Third, we consider the effect of a rise in the cost of disturbances to the principal. We show that all three changes increase the optimal intensity of intervention, but only the first also increases its likelihood and unambiguously decreases its duration.

In addition to the characterization of the efficient equilibrium of our model, we connect our theoretical framework and results to three historical episodes of indirect control: the Early Roman Empire, the Israeli-Palestinian conflict, and the Chechen wars. We describe how each of these situations is characterized by a government seeking to control a local authority who has the capacity to hinder disturbances in the presence of asymmetric information. We show that these situations feature the occasional use of military interventions by the government as punishment for disturbances. We also show that interventions are temporary, repeated, and are often deemed excessively destructive by outside observers. Finally, we connect some features of the examples to particular comparative statics of our model. Overall, these examples show how our model can be used to analyze and understand the use of military interventions in situations

of indirect control.

Related Literature

This paper is related to several different literatures. First, our paper contributes to the political economy literature on dynamic conflict by providing a formal framework for investigating the transitional dynamics between conflict and cooperation.⁵ The key distinction from the few related models which feature recurrent fighting (e.g., Anderlini, Gerardi, and Lagunoff, 2009, Fearon, 2004, Powell, 2009, and Rohner, Thoenig, and Zilibotti, 2010) is that we allow for levels of force which exceed the static best response, and we explicitly consider efficient equilibria. In doing so, we show that high levels of force are sustained by future cooperation, which allows for an analysis of the optimal intensity and duration of fighting. Because we focus on situations of indirect control in which one player uses violence to provide incentives to another player, our model bears a similar structure to Yared (2010). In contrast to this work, we introduce variable intervention intensity which allows for payoffs below the repeated static Nash equilibrium, and therefore optimal phases of intervention cannot last forever and must necessarily precede phases of cooperation.

Second, 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 principal lacks the commitment to punish.⁶ 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. As such, 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 new to our understanding of optimal punishments.⁷

Third, our paper contributes to the theoretical literature on the repeated moral hazard problem.⁸ A common feature of the baseline repeated moral hazard model is the absence of long run dynamics in the agent's continuation value.⁹ In contrast to this work, long run dynamics in the agent's continuation value emerge in our setting. Our model departs from the baseline in two respects. First, the structure of our stage game allows the principal to take over the agent's

⁵Some examples of work in this literature are Acemoglu and Robinson (2006), Baliga and Ely (2010), Chassang and Padró i Miquel (2009,2010), Jackson and Morelli (2008), and Powell (1999). Schwarz and Sonin (2008) 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.

⁶Some examples of models of punishments are Acemoglu and Wolitsky (2009), Chwe (1990), Dal Bó and Di Tella (2003), Dal Bó, Dal Bó and di Tella (2006), and Polinski and Shavell (1979,1984).

⁷Because 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.

⁸The literature on repeated moral hazard is vast and cannot be summarized here. Some examples are Ambrus and Egorov (2009), Phelan and Townsend (1991), Radner (1985), Rogerson (1985), and Spear and Srivastava (1987).

⁹This result is best elucidated in the continuous time model of Sannikov (2008) who shows that in the optimal contract, the principal backloads incentives to the agent, so that the agent is eventually fired or retired with a severance payment. The absence of long run dynamics in the agent's continuation value is also observed in discrete time models. This is true for instance in the model of DeMarzo and Fishman (2007) and in the computed examples of Phelan and Townsend (1991).

action at an endogenous cost to himself and the agent. Importantly, this departure on its own does *not* lead to any long run dynamics, as we show in the paper. Second, the principal suffers from limited commitment.¹⁰ We show that the combination of limited commitment together with the modified structure of the stage game generates the long run dynamics in the agent's continuation value. These dynamics emerge because of the alternating provision of incentives to the principal and to the agent.

Finally, note that the economics behind the dynamics of repeated intervention in our principal-agent framework are related to the insights in the literature on price wars under oligopolistic competition (e.g., Green and Porter, 1984).¹¹ As in our environment, this literature highlights how the equilibrium realization of statically inefficient outcomes (such as price wars) can serve to sustain cooperation. Because of the technical complexity associated with multi-sided private action, this literature imposes restrictions on players' strategies in order to characterize the dynamics of non-cooperation, making it difficult to address the tradeoffs underlying the optimal contract or the comparative statics.¹² In contrast to this framework, private information in our principal-agent environment is one-sided, which reduces this technical complexity and allows us to explicitly characterize the dynamics of non-cooperation together with the tradeoffs and comparative statics underlying the fully optimal contract with general history-dependent strategies.¹³

The paper is organized as follows. Section 2 describes the model. Section 3 defines the efficient sequential equilibrium. Section 4 characterizes the equilibrium under full commitment by the principal and highlights the absence of long run dynamics. Section 5 characterizes the equilibrium under limited commitment by the principal and provides our main results. Section 6 provides extensions, and we discuss our results in the context of some historical episodes in Section 7. Section 8 concludes. Appendix A contains the most important proofs and additional material not included in the text. Appendix B, which is available online, includes additional proofs not included in Appendix A.

¹⁰Other work introduces limited commitment on the side of the principal. For instance, Levin (2003) and Halac (2010) consider the role of an outside option for the principal and they find no long run dynamics for the agent's continuation value since, if the agent is not fired, then the long run contract corresponds to a stationary contract. Fong and Li (2010) also consider limited commitment on the side of the principal, though they show that long run dynamics could emerge under the assumption that low effort by the agent provides lower utility to the agent than termination, but this assumption does not hold in our setting.

¹¹For related work on price wars, see also Porter (1983), Radner (1986), Rotemberg and Saloner (1986), and Staiger and Wollack (1992). A different treatment dynamic oligopoly considers the role of private information (e.g., Athey, Bagwell, and Sanchirico, 2004) as opposed to private action, which is our focus.

¹²See Mailath and Samuelson (2006, p. 347-54) for an exposition of these difficulties. The work of Sannikov (2007) suggests some of these difficulties can be addressed in a continuous time framework.

¹³To do this, we utilize the recursive techniques developed by Abreu, Pearce, and Stacchetti (1990).

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, the principal can withhold force and allow the agent to exert *unobservable* effort in controlling disturbances. In this situation, if a large disturbance occurs, the principal cannot determine if it is due to the agent's negligence or due to bad luck. On the other hand, the principal can directly intervene to control disturbances himself, and in doing so he chooses the intensity of force. Both the principal and the agent suffer from limited commitment. In our benchmark model, we rule out transfers from the principal to the agent—which are standard in the dynamic principal-agent literature—since our focus is on the use of interventions. This is done purely for expositional simplicity. We allow for transfers in Section 6.1.1 and show that none of our results regarding the dynamics of intervention 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 $f_t = 0$, then the principal does not intervene and the agent *privately* chooses whether to exert high effort ($e_t = \eta$) or low effort ($e_t = 0 < \eta$) in minimizing disturbances. Nature then stochastically chooses the size $s_t \in S \equiv [0, \bar{s}]$ of a publicly observed disturbance. The principal receives $-s_t\chi$ from a disturbance where $\chi > 0$ parameterizes the cost of disturbances to the principal. Independently of the shock s_t , the agent loses e_t from exerting effort. The c.d.f. of s_t conditional on e_t is $\Phi(s_t, e_t)$. We let $\Phi(s_t, 0) < \Phi(s_t, \eta)$ for $s_t \in (0, \bar{s})$ so that higher disturbances are more likely under low effort. Therefore, letting $\pi_a(e_t)$ correspond to the expected value of s_t conditional on e_t , it follows that $\pi_a(\eta) < \pi_a(0)$ so that high effort reduces the expected size of a disturbance.¹⁴ The parameter η captures the cost of effort to the agent.¹⁵ We make the following technical assumptions on $\Phi(s_t, e_t)$. For all $s_t \in S$ and $e_t = \{0, \eta\}$, $\Phi(s_t, e_t) > 0$ and $\Phi(s_t, e_t)$ is twice continuously differentiable with respect to s_t . $\Phi(s_t, e_t)$ also satisfies the monotone likelihood ratio property (MLRP) so that $\Phi_s(s_t, 0)/\Phi_s(s_t, \eta)$ is strictly increasing in s_t , and we let $\lim_{s_t \rightarrow \bar{s}} \Phi_s(s_t, 0)/\Phi_s(s_t, \eta) = \infty$ so as to guarantee interior solutions. Finally, $\Phi_{ss}(s_t, e_t) < 0$, so that it is concave.¹⁶

If $f_t = 1$, then the principal publicly decides the intensity of force $i_t \in [0, \bar{i}]$. In this case,

¹⁴Due 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 large enough disturbance could stochastically force the principal to make a permanent concession beneficial to this player. Under this interpretation, the principal may be able to unilaterally make a concession to end all disturbances, a situation we consider in Section 6.1.2.

¹⁵The 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 large 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 large disturbances.

¹⁶While we model disturbances as continuous, one could alternatively model disturbances as a binary event with $s_t = \{0, 1\}$ without altering any of our main results. Details available upon request.

the payoff to the principal is $-\pi_p\chi - Ai_t$ and the payoff to the agent is $-g(i_t)$, where $A > 0$ and $g(i_t) > 0 \forall i_t \geq 0$. The parameter A captures the marginal cost of intensive force.¹⁷ Within the term $-\pi_p\chi - Ai_t$ is embedded the cost of a stochastic disturbance, where π_p represents the expected size of such a disturbance conditional on intervention. Analogously, within the term $-g(i_t)$ is the cost of the damage suffered by the agent when the principal intervenes, where this damage is increasing in intensity i_t .¹⁸ We let $g''(i_t) < 0$ with $\lim_{i_t \rightarrow 0} g'(i_t) = \infty$ and $\lim_{i_t \rightarrow \bar{i}} g'(i_t) = 0$. The concavity of $g(\cdot)$ captures the fact that there are diminishing returns to the use of intensity by the principal.

Importantly, conditional on intervention by the principal, both the principal and the agent are strictly better off under $i_t = 0$. This is because choosing $i_t > 0$ imposes more damage on the agent, it is costly to the principal, and it does not directly 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 proper interpretation of $i_t = 0$ is therefore not the absence of force, but rather the principal's statically optimal level of force, meaning the level of intensity associated with the principal seeking to directly minimize immediate disturbances. As an example, suppose that the principal was interested in limiting riots, and suppose that, given the costs, the statically optimal means of doing so for the principal is to impose a curfew only on the neighborhoods which are more riot-prone. In this situation, excessive force (i.e., $i_t > 0$) corresponds to imposing broader-based curfews in the region and engaging in other forms of harassment or destruction. These additional actions have minimal direct effect on reducing riots but they certainly impose additional costs on both the principal and the agent in the region.¹⁹

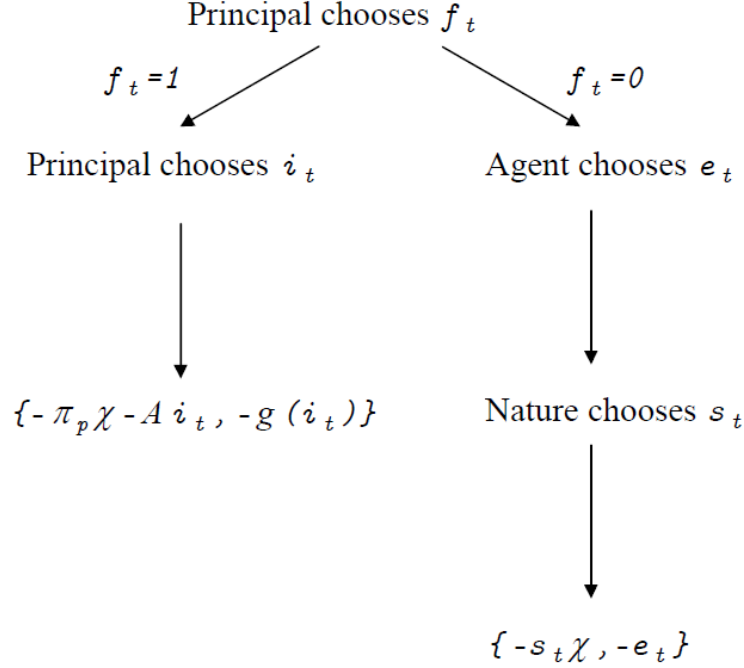
¹⁷For instance, A can decline if there is less international rebuke for the use of force.

¹⁸In 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.

¹⁹Thus, $-g(0)$ corresponds to the agent's disutility under the principal's statically optimal level intensity. This normalization has no qualitative effect on our results and yields considerable notational ease. In general, intensity could affect the size π_p of a disturbance under an intervention. If π_p is a convex function of intensity, one can show that the efficient sequential equilibrium only features levels of intensity above the statically optimal level. Details available upon request.

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 player $j = \{p, a\}$ at the beginning of the stage game 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 $-\eta > -g(0)$.

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 minimizing a disturbance. Intuitively, the agent is better informed about the sources of disturbances and is better than the principal at reducing them. Moreover, from an ex-ante perspective, the agent prefers to exert high effort to control disturbances versus enduring the damage from any intervention by the principal. In sum, this assumption implies that allowing the agent to exert effort dominates intervention by the principal.

Assumption 2 states that the principal is strictly better off using intervention to minimize disturbances versus letting the agent exert low effort. This assumption has an important implication. Specifically, in a one-shot version of this game, intervention with minimal intensity (i.e., $f_t = 1$ and $i_t = 0$) is the unique static Nash equilibrium. This is because conditional on no intervention (i.e., $f_t = 0$), the agent chooses minimal effort (i.e., $e_t = 0$). 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.²⁰ We refer to this situation with $f_t = 1$ and $i_t = 0$ as *direct control*.

Note that we have implicitly assumed that there is no asymmetry of information during intervention by the principal. There are two ways to interpret this assumption in our context. First, if the principal takes over the task—as occurs when he exerts direct control—the agent may have no reason to exert effort as he is made redundant. Second, during the disruptive and violent interventions that are the focus of our analysis, the agent may be sufficiently incapacitated that he cannot actually exert high effort. In both of these cases, asymmetric information during intervention is clearly less of a concern since the agent’s effort is largely irrelevant.²¹ This formulation, from a technical standpoint, also has the advantage of making the equilibrium tractable, since incentives need only be provided for one player in any given period.

Permanent direct control is always a sequential equilibrium of the repeated game. However, since it is inefficient (by Assumption 1), repeated game strategies may be able to enhance the welfare of both players. Nevertheless, there are three 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 large disturbance occurs, the principal cannot determine if this is accidental (i.e., $e_t = \eta$) or if this is intentional (i.e., $e_t = 0$).

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 Appendix A. 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.²²

Since we are concerned with optimal policy, we characterize the set of equilibria which

²⁰ Assumption 2 facilitates exposition by guaranteeing the existence of long run interventions. 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 since interventions must still occur in expectation. Section 6.1.2 provides an extension with a permanent concession which is isomorphic to this scenario.

²¹ Moreover, even if the agent’s effort under intervention did actually matter, the principal would presumably find it easier to monitor him under intervention.

²² Because 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).

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 credible punishment. More specifically, all public deviations from equilibrium actions by a given player lead to his worst credible punishment off the equilibrium path, which we denote by $\underline{J}$ for the principal and $\underline{U}$ for the agent. Note that

$$\begin{aligned}\underline{J} &= -\frac{\pi_p \chi}{1-\beta} \text{ and} \\ \underline{U} &\leq -\frac{g(0)}{1-\beta}.\end{aligned}$$

The principal cannot receive a lower payoff than under permanent direct control, as he could revert to it at any point. Moreover, the agent can be credibly punished by the principal at least as harshly as under permanent direct control, which is the repetition of the static Nash equilibrium. Therefore, while the structure of the game determines $\underline{J}$, $\underline{U}$ is determined endogenously by equilibrium strategies.

We allow players to choose correlated strategies so as to potentially randomize over the choice of intervention, intensity, and effort. We do this since this could improve efficiency. Formally, 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 uniform c.d.f. 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 choice.

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.²³ 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 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. Let 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 $z = \{z^1, z^2\}$. Let U_z^F represent 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 and the size of the disturbance s is $U_{z,s}^N$. 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 .

Associated with U is $J(U)$, which represents the highest continuation value achievable by

²³This is a consequence of the insights from the work of Abreu, Pearce, and Stacchetti (1990).

the principal in a sequential equilibrium conditional on the agent achieving a continuation value of U . Letting $\rho = \left\{ f_z, i_z, e_z, U_z^F, \{U_{z,s}^N\}_{s \in [0, \bar{s}]} \right\}_{z \in Z}$, the recursive program which characterizes the efficient sequential equilibrium is

$$J(U) = \max_{\rho} E_z \left\{ f_z \left(-\pi_p \chi - A i_z + \beta J(U_z^F) \right) + (1 - f_z) \left(-\pi_a(e_z) \chi + \beta E_s \left\{ J(U_{z,s}^N) | e_z \right\} \right) \right\} \quad (1)$$

s.t.

$$U = E_z \left\{ f_z \left(-g(i_z) + \beta U_z^F \right) + (1 - f_z) \left(-e_z + \beta E_s \left\{ U_{z,s}^N | e_z \right\} \right) \right\}, \quad (2)$$

$$-\pi_p \chi - A i_z + \beta J(U_z^F) \geq \underline{J} \quad \forall z^1, z^2 \quad (3)$$

$$E_z \left\{ -\pi_a(e_z) \chi + \beta E_s \left\{ J(U_{z,s}^N) | e_z \right\} | z^1 \right\} \geq \underline{J} \quad \forall z^1 \quad (4)$$

$$\beta \left(E_s \left\{ U_{z,s}^N | e_z = \eta \right\} - E_s \left\{ U_{z,s}^N | e_z = 0 \right\} \right) \geq e_z \quad \forall z^1, z^2 \quad (5)$$

$$J(U_z^F), J(U_{z,s}^N) \geq \underline{J} \quad \forall z^1, z^2, s \quad (6)$$

$$U_z^F, U_{z,s}^N \geq \underline{U} \quad \forall z^1, z^2, s \quad (7)$$

$$f_z \in [0, 1], i_z \geq 0, \text{ and } e_z = \{0, \eta\} \quad \forall z^1, z^2. \quad (8)$$

(1) represents the continuation value to the principal written in a recursive fashion at a given history. Equation (2) represents the promise keeping constraint which ensures that the agent is achieving a continuation value of U . Constraints (8) ensure that the allocation is feasible. Constraints (3) – (7) 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) – (7) capture the inefficiencies introduced by the presence of limited commitment and imperfect information which ultimately lead to the need for intervention. Constraints (3) and (4) take into account that at any history, the principal cannot commit to refraining from permanent direct control which provides a continuation value of $\underline{J}$. Specifically, constraint (3) captures the fact that at any history, the principal cannot commit to using intensive force since this is costly. It ensures that at any history in which the principal intervenes with $i_z > 0$, the principal is rewarded for this in the future with $J(U_z^F) > \underline{J}$. Constraint (4) captures the fact that at any history, the principal may not be able to commit to allowing the agent to exert low effort, and if that is the case, the principal is rewarded for this in the future with $E_s \left\{ J(U_{z,s}^N) | e_z \right\} > \underline{J}$.²⁴ Constraint (5) 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 (5) ensures that the agent's punishment from this deviation is weakly exceeded by the equilibrium path reward for choosing high effort.²⁵ Constraints (6) and (7) guarantee that the set of continuation values

²⁴More precisely, following the realization of z^1 , it must be that the expectation of effort by the agent—which can depend on the realization of z^2 —is sufficiently high or alternatively that the continuation value to the principal is sufficiently high.

²⁵Note that we have ignored the constraint that the agent does not deviate to high effort if $e_z = 0$. We explicitly

$\left\{U_z^F, \{U_{z,s}^N\}_{s \in [0, \bar{s}]}\right\}_{z \in Z}$ chosen in the future satisfy future incentive compatibility constraints for the principal and for the agent, where $\underline{U}$ corresponds to the lowest continuation value for the agent.

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

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

This definition states that the intensity of intervention is the expected intensity of the force used by the principal during intervention; the likelihood of intervention is the probability that the principal intervenes following a period of non-intervention; and the duration of intervention is the probability that intervention continues into the next period.²⁶ To facilitate exposition, we assume that players are sufficiently patient for the remainder of our discussion.

Assumption 3 (High Patience) $\beta > \hat{\beta}$ for some $\hat{\beta} \in (0, 1)$.

The exact value of $\hat{\beta}$ is described in detail in Appendix B.²⁷

4 Full Commitment Benchmark

In this section, we characterize the equilibrium in the presence of full commitment by the principal (i.e., in the model which ignores constraints (3), (4), and (6)) in order to show that no intervention dynamics emerge in this setting.

To this end, it is useful to note that the agent cannot possibly receive a continuation value which exceeds 0, since this is the continuation value associated with the agent providing low effort forever under no intervention by the principal. Therefore, $U_t \leq 0$ for all t . The following proposition characterizes the dynamics of intervention in the efficient sequential equilibrium under full commitment by the principal.

Proposition 1 (full commitment) *Under full commitment by the principal, the following is true for all t :*

1. *If $U_t < 0$, then the probability of intervention in the future is positive, so that $\Pr\{f_{t+k} = 1\} > 0$ for some $k \geq 0$,*

consider this constraint in Appendix A and we can show that it never binds in equilibrium.

²⁶The expected length of time that $f_t = 1$ is equal to $1/(1 - \Pr\{f_{t+1} = 1 | f_t = 1\})$ which is a monotonic transformation of $\Pr\{f_{t+1} = 1 | f_t = 1\}$, and for this reason we interpret $\Pr\{f_{t+1} = 1 | f_t = 1\}$ as corresponding to the duration of intervention.

²⁷This assumption allows us to explicitly characterize the equilibrium since it implies that the threat of intervention with minimal force is sufficient to induce high effort and that the equilibrium duration of intervention is bounded away from 0. We can show that for any discount factor $\beta \in (0, 1)$, there exists a set of parameter values for which $\hat{\beta}$ is sufficiently low that $\beta > \hat{\beta}$ so that Assumption 3 holds. Details available upon request.

2. *The duration of intervention $\Pr \{f_{t+1} = 1 | f_t = 1\} = 1$ so that intervention is permanent once it is used, and*
3. *The intensity of intervention $E \{i_t | f_t = 1\}$ is positive and it is constant so that conditional on $f_t = 1$, $i_{t+k} = i_t$ for all $k \geq 0$.*

The first part of Proposition 1 states that if the agent is receiving a continuation value below that associated with contributing zero effort forever, then there is a positive probability of intervention in the future. The argument behind this result is straightforward. Suppose by contradiction that the probability of intervention going forward were zero. Then the agent has no incentive to exert high effort and could choose to exert low effort forever without detection and achieve a continuation value equal to 0, which would make him strictly better off since $U_t < 0$. Intuitively, if $U_t < 0$, then interventions must take place in the future in order to induce the agent to exert high effort in minimizing disturbances. Therefore, even though interventions are ex-post inefficient by Assumption 1, they are ex-ante efficient since they provide the right incentives to the agent to exert effort.²⁸

The second and third parts of Proposition 1 state that there are no dynamics of intervention since interventions last forever and entail a fixed level of intensity. The intuition for the second part is that permanent interventions are the most efficient means of providing ex-ante incentives for high effort by the agent. Formally, suppose by contradiction that there is no intervention at t and, with some probability, the principal intervenes starting from $t + 1$ for k periods. The principal could easily choose an alternative policy of intervening forever but with a lower likelihood (i.e., being more forgiving of small disturbances). This change in policy can be done in a way that does not change his own welfare or that of the agent's at t . Note that because of the MLRP property, this change in policy is also better for the agent's incentives and relaxes his incentive constraint (5). This is because the punishment for very high disturbances—which are more likely under low effort—becomes more severe. Importantly, since the agent's incentive constraint is slackened, and since intervention at $t + 1$ is costly to both players, both the principal and the agent can be made strictly better off at t by reducing the likelihood of intervention even further. Consequently, lengthening interventions is good for ex-ante welfare for both parties. The third part of the proposition follows from the fact that changing the level of intensity during an intervention cannot improve the agent's ex-ante incentives for putting in high effort since $g(\cdot)$ is concave. Moreover, since the marginal cost of intensity for the principal is constant, he cannot improve his own welfare conditional on permanently intervening by choosing a non-constant intensity.²⁹

Even though we do not focus on equilibrium dynamics outside of phases of intervention, it is useful to briefly describe them in order to understand the mechanics of the model. In

²⁸More generally, efficient sequential equilibria both under full commitment and limited commitment by the principal are not renegotiation proof. According to the definition of Farrell and Maskin (1989), the only weakly renegotiation proof equilibrium in our setting is the repeated static Nash equilibrium.

²⁹In Section 5.4, we discuss the range of intensity which is chosen by the principal.

providing the agent with a continuation value U_t , the principal decides whether to intervene, and if he does not intervene, whether to request high or low effort from the agent. It is clear that if the continuation value U_t is low enough, the principal punishes the agent by intervening, and Proposition 1 implies that intervention is permanent. Analogously, if U_t is high enough, the principal rewards the agent by not intervening and allowing him to exert low effort. In the intermediate range of continuation values, the principal does not intervene and requests high effort from the agent. Incentive provision for the agent in this range implies that small disturbances are rewarded with an increase in continuation value (and hence a higher likelihood of non-intervention and low effort in the future) and large disturbances are punished with a decrease in continuation value (and hence a higher likelihood of intervention in the future). Therefore, to provide incentives to the agent, the principal chooses a level of intensity that he commits to exerting forever in the event that a sufficient number of large disturbances occur along the equilibrium path. If instead small disturbances occur along the equilibrium path, then the principal rewards the agent by allowing him to exert low effort. Eventually, such a well-performing agent can exert low effort forever. The implication of these equilibrium dynamics is described in the below corollary.

Corollary 1 *Under full commitment by the principal, the long run equilibrium must feature one of the two following possibilities:*

1. *Permanent intervention by the principal, or*
2. *Permanent non-intervention by the principal with low effort by the agent.*

This corollary highlights the importance of full commitment by the principal to the long run contract between the principal and the agent. Both long run outcomes provide the principal with a strictly lower payoff than what he can guarantee himself through permanent direct control (i.e., intervention with minimal intensity). This means that if given the option ex-post, the principal would choose to change the terms of this dynamic contract.

5 Equilibrium under Limited Commitment

In this section, we characterize the equilibrium in our environment, which takes into account limited commitment by the principal. In Section 5.1, we show that repeated temporary interventions must occur. In Section 5.2, we consider the optimal intensity, likelihood, and duration of intervention, and we characterize an important tradeoff in the optimal contract between the duration and the intensity of intervention. In Section 5.3 we consider comparative statics. Finally, in Section 5.4, we discuss the distortions which emerge as a consequence of the principal's inability to commit.

5.1 Repeated Intervention

The previous section shows that in the presence of full commitment by the principal, there are no intervention dynamics. We now consider the solution to the full problem in (1) – (8) which takes into account the principal’s inability to commit.

Proposition 2 (*limited commitment*) *Under limited commitment by the principal, the following is true for all t :*

1. *The probability of intervention in the future is always positive, so that $\Pr \{f_{t+k} = 1\} > 0$ for some $k \geq 0$,*
2. *The duration of intervention $\Pr \{f_{t+1} = 1 | f_t = 1\}$ satisfies $\Pr \{f_{t+k} = 1 | f_t = 1\} < 1$ for some $k \geq 0$ so that intervention is temporary, and the agent never exerts low effort following intervention so that $\Pr \{e_{t+k} = 0 | f_t = 1\} = 0$ for all $k \geq 0$, and*
3. *The intensity of intervention $E \{i_t | f_t = 1\}$ is positive and equal to a unique constant i^* which satisfies*

$$\frac{A}{g'(i^*)} = \frac{(\pi_p - \pi_a(\eta))\chi + Ai^*}{g(i^*) - \eta}. \quad (9)$$

Note the differences between each parts of Proposition 2 under limited commitment and Proposition 1 under full commitment. The first part of Proposition 2 states that future interventions always occur with positive probability, whereas the first part of Proposition 1 states that this is only true if the continuation value to the agent today is below 0. The second part of Proposition 2 states that interventions are always temporary and followed by high effort by the agent, whereas the second part of Proposition 1 states that interventions are always permanent. Finally, the third parts of both propositions state that a constant intensity is used during intervention, and Proposition 2 explicitly characterizes this level of intensity in the case of limited commitment.

The reasoning behind each part of Proposition 2 is as follows. To understand the first part, note that if $U_t < 0$ there is always a positive probability of intervention in the future, and this follows from the same arguments as under the full commitment case described in the discussion of Proposition 1. The reason why this is true always here is that it is not possible for $U_t = 0$ since in this situation the agent exerts low effort forever, and by Assumption 2, this is not incentive compatible for the principal since he would prefer to intervene.³⁰

The logic behind the second part of Proposition 2 follows from the presence of limited commitment on the side of the principal. Permanent intervention with positive intensity as under full commitment is not credible since the principal would prefer to deviate from such an arrangement by intervening with minimal intensity and receiving the repetition of his static Nash payoff $\underline{J}$. Consequently, periods of non-intervention in which the agent exerts high effort

³⁰This result is also present in Yared (2010).

must occur in the future in order to reward a principal who is intervening today. Moreover, the proposition also states that the agent *never* exerts low effort following an intervention. The reason is that this is optimal for the provision of ex-ante incentives to the agent. Recall from the discussion in Section 4 that providing optimal incentives requires maximizing the duration of intervention in the future, since this maximally relaxes the agent's incentive compatibility constraint along the equilibrium path. Since permanent intervention is not credible given the incentive compatibility constraint of the principal, an additional means of punishing the agent following a sequence of large disturbances is to ask for maximal effort by the agent during periods of non-intervention.

Even though we do not focus on equilibrium dynamics outside of phases of intervention, it is useful to briefly describe them. If the continuation value U_t is low enough, the principal punishes the agent by intervening. If U_t is high enough, the principal rewards the agent by not intervening and allowing him to exert low effort temporarily. In the intermediate range of U_t , the principal does not intervene and requests high effort by the agent, where small disturbances are rewarded with an increase in continuation value and large disturbances are punished with a decrease in continuation value. The implication of these equilibrium dynamics is described in the below corollary.

Corollary 2 *Under limited commitment by the principal, the long run equilibrium must feature fluctuations between periods of intervention with a constant intensity i^* and periods of non-intervention with high effort by the agent.*

This corollary and our above discussion highlight some similarities and some differences between this environment and the case of full commitment by the principal. As in the case of full commitment, high effort by the agent in the intermediate range of continuation values is followed by an increase or decrease in continuation value, depending on the size of the disturbance. However, in contrast to the full commitment case, there are long run equilibrium dynamics in the continuation value to the agent which emerge here. In particular, the agent is never punished with permanent intervention or rewarded with permanent non-intervention and low effort, as in the case of full commitment.

The last part of Proposition 2 states that the level of intensity used during intervention is positive and equal to a unique constant which satisfies (9). This level of intensity is constant over phases of intervention, as in the case of full commitment, and it is described in more detail in the following section.

5.2 Intensity, Likelihood, and Duration of Intervention

In this section, we more thoroughly investigate the intensity, likelihood, and duration of intervention which emerge in the efficient equilibrium. The second part of Proposition 2 shows that once intervention has occurred, the principal and the agent alternate between intervention with

intensity i^* and non-intervention with high effort by the agent. Note that characterization of the equilibrium under intervention is complicated by the fact that, even though the intensity of intervention is unique, the likelihood and duration of intervention are not unique since there are multiple ways of satisfying the principal's incentive compatibility constraint. For instance, the duration of intervention can be high if the principal expects a lengthy period of non-intervention in the future (i.e., a low likelihood of intervention in the future). Alternatively, the duration of intervention can be low if the principal expects a short period of non-intervention in the future (i.e., a high likelihood of intervention in the future). This multiplicity of potential solutions does not emerge in the case of full commitment since intervention lasts forever.

To alleviate this multiplicity and further characterize the likelihood and duration of intervention, we focus on the solution which satisfies the *Bang-Bang* property as described by Abreu, Pearce, and Stacchetti (1990). 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. One can show that this leads us to select the equilibrium with the highest intervention duration, which allows us to facilitate comparison with the full commitment case.³¹ The next proposition as well as all the remaining results of the paper characterize the equilibrium which satisfies the *Bang-Bang* property.

Proposition 3 (*characterization of Bang-Bang equilibrium*) *Under limited commitment by the principal, if intervention has occurred before t (i.e., $f_{t-k} = 1$ for some $k \geq 0$), then the equilibrium at t features either a cooperative phase or a punishment phase which are characterized as follows:*

1. *In the cooperative phase at t , the principal does not intervene and the agent exerts high effort (i.e., $f_t = 0$ and $e_t = \eta$). If the size of the disturbance is below a threshold (i.e. $s_t < \tilde{s}(i^*) \in (0, \bar{s})$), then the cooperative phase restarts at $t + 1$, otherwise the players transition to the punishment phase at $t + 1$, and*
2. *In the punishment phase at t , the principal intervenes with intensity i^* (i.e., $f_t = 1$ and $i_t = i^*$). With probability $\tilde{d}(i^*) \in (0, 1)$, the punishment phase restarts at $t + 1$, and with probability $1 - \tilde{d}(i^*)$ the players transition to the cooperative phase at $t + 1$.*

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 large enough 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

³¹This equilibrium also corresponds to the unique solution if we impose an additional constraint that players only have one period memory. Details available upon request.

trigger the agent to choose low effort in all future cooperative phases, making direct control—i.e., permanent intervention with minimal intensity—a necessity.³²

To see the mechanics of this equilibrium, let $\overline{U}(i)$ and $\underline{U}(i)$ correspond to the agent's continuation value during cooperation and punishment, respectively given the level of intensity $i = i^*$. These satisfy

$$\overline{U}(i) = -\eta + \beta (\Phi(\tilde{s}(i), \eta) \overline{U}(i) + (1 - \Phi(\tilde{s}(i), \eta)) \underline{U}(i)), \text{ and} \quad (10)$$

$$\underline{U}(i) = -g(i) + \beta \left((1 - \tilde{d}(i)) \overline{U}(i) + \tilde{d}(i) \underline{U}(i) \right). \quad (11)$$

(10) shows that in the cooperative phase, the agent exerts high effort today and faces two possibilities tomorrow. If a large enough disturbance occurs, play moves to punishment and he obtains $\underline{U}(i)$. Otherwise, cooperation is maintained and he receives $\overline{U}(i)$ tomorrow. Since the agent is punished for disturbances which exceed $\tilde{s}(i)$, it follows that the likelihood of punishment $\tilde{l}(i)$ satisfies:

$$\tilde{l}(i) = \Pr \{f_{t+1} = 1 | f_t = 0\} = 1 - \Phi(\tilde{s}(i), \eta). \quad (12)$$

(11) shows that in the punishment phase, the agent endures punishment with intensity $i = i^*$ today, and he receives $\underline{U}(i)$ tomorrow with probability $\tilde{d}(i)$ and $\overline{U}(i)$ tomorrow with probability $1 - \tilde{d}(i)$.³³

According to (10), the agent chooses high effort during the cooperative phase, which means his incentive compatibility constraint (5) must be satisfied. Moreover, one can show that it binds so that:

$$\beta (\Phi(\tilde{s}(i), \eta) - \Phi(\tilde{s}(i), 0)) (\overline{U}(i) - \underline{U}(i)) = \eta. \quad (13)$$

Let $\overline{J}(i)$ and $\underline{J}(i)$ correspond to the principal's continuation value during cooperation and punishment, respectively, given the level of intensity $i = i^*$. These satisfy

$$\overline{J}(i) = -\pi_a(\eta) \chi + \beta (\Phi(\tilde{s}(i), \eta) \overline{J}(i) + (1 - \Phi(\tilde{s}(i), \eta)) \underline{J}(i)), \text{ and} \quad (14)$$

$$\underline{J}(i) = -\pi_p \chi - Ai + \beta \left((1 - \tilde{d}(i)) \overline{J}(i) + \tilde{d}(i) \underline{J}(i) \right). \quad (15)$$

(14) shows that during cooperation the principal suffers from disturbances with expected size $\pi_a(\eta)$, and (15) shows that during punishment the principal suffers from disturbances with

³²More precisely, permanent direct control is one of many means of implementing the worst punishment for the principal. There are many other continuation games which provide the principal with a continuation value of $\underline{J}$ which can serve as punishment. For instance, the principal and the agent can transition to the equilibrium which provides the agent with the highest credible continuation value.

³³The principal uses the public signal z_t^1 to randomly exit from the punishment phase. In practice, one can interpret this public signal as corresponding to the random size of disturbances which occur under intervention which are visible to both the principal and the agent. Alternatively, one can imagine that with fine enough time periods, the principal can intervene for a fixed deterministic interval of time—which is clearly observable to both players—before resuming cooperation.

³⁴Letting (5) bind is optimal since it minimizes the likelihood of punishment which is costly to both the principal and the agent.

expected size π_p and he also suffers from intervening with intensive force.

Note that during punishment the principal weakly prefers choosing intensive force—which is statically dominated—and being rewarded for it with cooperation in the future versus choosing his optimal deviation which is to intervene forever with minimal intensity. One can show that the incentive compatibility constraint on the principal (3) binds so that

$$\underline{J}(i) = \underline{J}. \quad (16)$$

Given i^* , note that (10) – (16) forms a system of equations which allows us to trace exactly how the cooperative and punishment phases are tightly linked. In particular, the value of $\underline{U}(i)$ sustains the values of $\bar{J}(i)$ and $\bar{U}(i)$. Equation (13) implies that, holding $\bar{U}(i)$ constant, the lower is $\underline{U}(i)$, the higher is $\tilde{s}(i)$, so that the more forgiving the principal can be towards the agent.³⁵ Intuitively, the harsher the punishment, the less often it needs to be used and the lower the likelihood of intervention. Moreover, the lower the likelihood of punishment, the better off are the principal and the agent during the cooperative phase. This is clear since, holding all else fixed, the right hand sides of (10) and (14) both increase as the likelihood of punishment $\tilde{l}(i)$ declines. As a consequence, the highest possible $\bar{J}(i)$ and $\bar{U}(i)$ are associated with the lowest $\underline{U}(i)$, as this makes for the longest sustainable cooperative phase—i.e., the lowest sustainable likelihood of punishment $\tilde{l}(i)$.

Similarly, the value of $\bar{J}(i)$ sustains the value of $\underline{U}(i)$. Equations (15) and (16) imply that, holding all else fixed, the higher is $\bar{J}(i)$, then the higher is the implied value of $\tilde{d}(i)$. This is because the higher the principal's continuation value under cooperation, the more easily can the principal be induced to punish for longer, as his value under punishment is anchored at $\underline{J}$. Moreover, if punishment is longer, the worse off is the agent during the punishment phase. This is clear since, holding all else fixed, the right hand side of (11) decreases as the duration of punishment $\tilde{d}(i)$ increases. As a consequence, the lowest possible $\underline{U}(i)$ is associated with the highest possible $\bar{J}(i)$, as this makes for the harshest punishment phase.

In order to see how these insights are related to the optimal level of intensity i^* , we can use equations (10) – (16) in order to consider how the equilibrium would differ under some alternative value of intensity $i \neq i^*$. More specifically, (10) – (16) corresponds to a system of seven equations and seven unknowns:

$$\left\{ \bar{U}(i), \underline{U}(i), \bar{J}(i), \underline{J}(i), \tilde{s}(i), \tilde{l}(i), \tilde{d}(i) \right\}. \quad (17)$$

Each unknown is a continuously differentiable function of i for $i \in [0, \bar{i}]$ for some $\bar{i} \in (i^*, \bar{i}]$, where $\bar{J}(i^*) = J(\bar{U}(i^*))$, and such a system of equations satisfies $\tilde{d}(i) \in [0, 1]$ and $\tilde{l}(i) \in [0, 1]$.³⁶

³⁵This follows from the fact that $\Phi(\tilde{s}(i), \eta) - \Phi(\tilde{s}(i), 0)$ is declining in $\tilde{s}(i)$ in the optimal contract since this minimizes the realization of punishment.

³⁶Formally, (10) – (16) admits two solutions for $\tilde{s}(i)$ and we select the higher value of $\tilde{s}(i)$, since this coincides with the solution in Proposition 3.

This means that it corresponds to a sequential equilibrium which has the same structure as the equilibrium described in Proposition 3 for some level of intensity i .³⁷ We can use the functions in (17) in order to better characterize the optimal level of intensity i^* .

Proposition 4 (*optimal intensity*) *In the system defined by (10) – (16), as intensity i increases from 0 to $\bar{i}$,*

1. *The duration of intervention $\tilde{d}(i) \in [0, 1]$ decreases,*
2. *The agent's continuation value of punishment $\underline{U}(i)$ decreases (increases) if intensity is below (above) i^* ,*
3. *The likelihood of intervention $\tilde{l}(i) \in [0, 1]$ decreases (increases) if intensity is below (above) i^* , and*
4. *The principal's and the agent's continuation value of cooperation, $\bar{J}(i)$ and $\bar{U}(i)$ increase (decrease) if intensity is below (above) i^* .*

Proposition 4 implies that the optimal level of intensity minimizes the agent's continuation value of punishment, minimizes the likelihood of intervention, and maximizes the principal's and the agent's continuation value of cooperation. These insights allow us to interpret for the value of i^* defined in (9).

The principal's incentives to intervene are the driving force behind Proposition 4. 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 (15) and (16) and it implies that $\tilde{d}'(i) < 0$, so that the duration of intervention is declining in intensity.

Now consider what this implies for the continuation value of the agent under punishment, $\underline{U}(i)$. At low levels of i , an increase in intensity naturally means that the prospect of punishment is worse for the agent, and $\underline{U}(i)$ 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 (15) and (16). As a consequence, above a certain i , $\underline{U}(i)$ becomes *increasing* in i .

Since the agent's continuation value under punishment first decreases and then increases with intensity, the likelihood of intervention $\tilde{l}(i)$ first decreases and then increases with intensity, as implied by (13). As the punishment for the agent becomes worse, a smaller likelihood of punishment is needed to satisfy his incentive constraint. As previously discussed, lower likelihood is better from the perspective of the principal and the agent because it maximizes the duration

³⁷More specifically, in such an equilibrium, (13) implies that the agent's incentive compatibility constraint binds during the cooperative phase, and (16) implies that the principal's incentive compatibility constraint binds during the punishment phase.

of cooperation (i.e., the probability of transitioning to the cooperative phase tomorrow starting from the cooperative phase today is maximized). Therefore, the principal always chooses the level of intensity that minimizes the likelihood of punishment. As stated in Proposition 4, this level is i^* .

We can now interpret the value of i^* as defined in (9). Note that the principal's objective is to minimize the agent's value under punishment, and he can do so by either changing the intensity or the duration of intervention. On the margin, increasing intensity reduces the agent's continuation value by $g'(i^*)$ and costs the principal A . Alternatively, increasing the duration of intervention by one period costs each player the difference in flow payoffs between intervention and non-intervention, so it reduces the agent's continuation value by $g(i^*) - \eta$ and costs the principal $(\pi_p - \pi_a(\eta))\chi + Ai^*$. Equation (9) shows that the optimal level of intensity used by the principal equalizes the marginal gains relative to the marginal costs of these two strategies for punishing the agent.

The fact that i^* is positive 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 sustained in the long run and the principal resorts to permanent direct control, as in Yared (2010). Intuitively, the punishment to the agent is not sufficiently dire to warrant its use by the principal. In this situation, the second part of Proposition 2 would not hold since intervention is necessarily permanent once it is used. This highlights how it is the combination of limited commitment by the principal together with the equilibrium use of costly interventions which generates repeated equilibrium interventions.

5.3 Comparative Statics

In this section, we consider the effect of three factors on the optimal intensity, likelihood, and duration of intervention. First, 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 for the agent to antagonize rival factions contributing to disturbances, or alternatively if he acquires a higher preference for the realization of disturbances.³⁸ Second, we consider the effect of a decline in the cost of intensity to the principal (A). Third, we consider the effect of a rise in the cost of disturbances to the principal (χ).³⁹ The comparative statics are summarized in the below proposition, where as a reminder, $\tilde{l}(i^*)$ and $\tilde{d}(i^*)$ defined in Section 5.2, correspond to the likelihood and duration of intervention, respectively.⁴⁰

Proposition 5 (*comparative statics*)

³⁸That is, holding all else fixed so that $\Phi(s, \eta)$ and $\Phi(s, 0)$ are unchanged.

³⁹One 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.

⁴⁰Note that the values of $l(i)$ and $d(i)$ depend *directly* on parameters of the model such as η , A , and χ . This means that $l(i^*)$ and $d(i^*)$ can be affected by η , A , and χ in an independent manner from the effect through the change in i^* .

1. *If the cost of effort η increases (decreases), then the intensity of intervention i^* increases (decreases), the likelihood of intervention $\tilde{l}(i^*)$ increases (decreases), and the duration of intervention $\tilde{d}(i^*)$ decreases (increases), and*
2. *If the cost of intensity A decreases (increases) or if the cost of disturbances χ increases (decreases), then the intensity of intervention i^* increases and the likelihood of intervention $\tilde{l}(i^*)$ decreases (increases).*

This proposition states that all three changes increase the optimal intensity of intervention. To see why intensity must rise, recall from Section 5.2 that the optimal level of intensity minimizes the agent's value under punishment. Since higher levels of intensity are associated with a shorter duration of punishment, the optimal level of intensity equalizes the marginal gain relative to the marginal cost to the principal of using either higher intensity or higher duration to punish the agent.

Now consider the first case where the cost of effort for the agent rises. This means that following punishment, the agent's payoff from exerting high effort is lower. Consequently, for a principal seeking to minimize the agent's continuation value of punishment, this means that the cost of reducing duration relative to the benefit of increasing intensity is diminished, so that higher intensity becomes optimal. Now suppose that the cost of intensity declines as in the second case. Then the principal increases intensity since it becomes more efficient to use intensity relative to duration in punishing the agent. Finally, suppose the cost of disturbances rises as in the third case. In this situation, the cost to the principal of a longer duration of punishment is increased, since the relative cost of not delegating to the agent rises. Consequently, the optimal policy is to increase the intensity of intervention.

Even though all three changes increase the optimal intensity of intervention, only the first also raises its likelihood. Specifically, 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 large enough disturbances. In contrast, 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.⁴¹

Now consider the effect of these changes on the duration of intervention. While an increase in the cost of effort to the agent reduces the duration of intervention, the other two changes have an ambiguous effect on the duration of intervention. This ambiguity is driven by the fact that the principal responds optimally to changes in the environment by increasing the level of intensity in all cases.

To see this, 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

⁴¹Note that if instead the cost of intensity were parameterized by some binding upper bound i'' on the level of intensity, then the optimal contract would feature $i^* = i''$. In this situation, the likelihood and duration of intervention decrease if i'' increases, and this follows from the results in Proposition 4.

5.2. Suppose the cost of effort to the agent rises but i does not change. 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. Now take into account the fact that the principal increases the level of intensity so as to mitigate the rise in the likelihood of intervention. Again, this increase in the intensity of intervention must entail an even further decline in the duration of intervention so as to preserve the incentives for the principal to intervene. This explains why in this case, the duration of intervention unambiguously declines.

Now consider the effect on duration 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.⁴² Now take into account the fact that the principal responds to the change in the environment by increasing the level of intensity. This increase in the intensity of intervention puts downward pressure on duration of intervention so as to preserve the incentives for the principal to intervene. Which of these forces on the duration dominates is ambiguous and depends on the exact value of parameters as well as the functional forms of $g(\cdot)$ and $\Phi(\cdot)$ which ultimately determine how much i^* changes in response to changes in the environment.⁴³

5.4 Distortions from Limited Commitment

We have shown that limited commitment on the side of the principal implies the presence of repeated equilibrium interventions. In contrast, a principal with full commitment power would prefer to intervene permanently. Therefore, it is clear that both the principal and the agent could be made strictly better off if the principal could commit to his actions. Intuitively, the principal's ability to commit to an extreme punishment for the agent slackens the agent's incentive compatibility constraint, making it possible to delay costly interventions.⁴⁴

To get a sense of how the principal's ability to commit to permanent intervention raises efficiency, consider the system of equations represented by (10) – (16) where we replace (16)

⁴²Formally, this is equivalent to stating that $d(i)$ is decreasing in A and increasing in χ for each given i .

⁴³As an aside, note that a natural question concerns how these changes in the environment affect the continuation values to the principal and to the agent during the cooperative phase, $\bar{J}(i^*)$ and $\bar{U}(i^*)$, respectively. It can be shown that an increase in η reduces both of these, that a decrease in A raises both of these, and that an increase in χ decreases $\bar{J}(i^*)$ but raises $\bar{U}(i^*)$. Details available upon request.

⁴⁴A similar intuition holds regarding how the principal rewards the agent by allowing periods of non-intervention and low effort. Since a principal with commitment can do this permanently, this also slackens the agent's incentive compatibility constraint.

with

$$J(\underline{U}(i)) = \underline{J} - \Delta. \quad (18)$$

$\Delta \geq 0$ parameterizes the strength of the principal's commitment, since, the higher is Δ , the lower is the principal's min-max, and the more slack is the principal's incentive compatibility constraint. It can be shown that for a given Δ which is not too large, an analogous result to Proposition 4 holds, so that the same tradeoffs apply in the determination of the optimal level of intensity.⁴⁵ Moreover, given these tradeoffs, the same value of i^* emerges as the value which minimizes the continuation value of punishment to the agent and minimizes the likelihood of intervention. This is because equation (9) makes clear that i^* is chosen so that the principal equalizes the marginal gains relative to the marginal costs of using higher intensity versus higher duration of intervention in punishing the agent. These considerations are not affected by the value of Δ .

Given (10) – (15) and (18), as Δ rises, the principal is able to intervene for longer under i^* so that the duration of punishment rises. As such, the continuation value under punishment for the agent decreases. Satisfaction of the incentive constraint for the agent means that the likelihood of punishment during cooperation declines, and this means that the continuation value to the principal and to the agent under cooperation rises. Eventually, Δ rises to a point where duration equals 1 so that the principal and the agent receive a continuation value associated with permanent intervention during the punishment phase. In this situation, the continuation values to the principal and to the agent during the cooperative phase are maximized and are much higher relative to the case of no commitment since the principal can essentially commit to a permanent intervention.

A natural question concerns how the level of intensity under lack of commitment compares to that under full commitment. We can show that a principal who can fully commit always chooses a level of intensity at least as high as a principal constrained by limited commitment.

Lemma 1 *Under full commitment, the principal chooses a level of intensity at least as high as i^* .*

Recall that the level of i^* defined in (9) equalizes the marginal gain relative to the marginal cost to the principal of using either higher intensity or higher duration to punish the agent. By Proposition 1, the principal with full commitment intervenes permanently. Intervening permanently with a level of intensity below i^* is inefficient since it is dominated by intervening with i^* for some limited duration, where this follows from (9). Hence, a principal with full commitment never intervenes with intensity below i^* . In contrast, a principal with full commitment power could choose a level of intensity above i^* , even though this is strictly dominated for a principal without the ability to commit. This is because such a level of intensity could only possibly be

⁴⁵It is nevertheless no longer necessarily true that the implied duration and likelihood of intervention are between 0 and 1 for $i \leq i^*$. The value of Δ cannot be so high that the duration is larger than 1.

efficient if applied permanently.⁴⁶

Note that as the agent becomes more patient, the same future threat of punishment will more easily induce effort. One may conjecture that this implies that the distortions due to the principal's inability to commit to a permanent intervention are reduced as players become sufficiently patient. This is stated formally in the following proposition.

Proposition 6 (*high discount factor*) *As β approaches 1,*

1. *The principal's and agent's flow payoffs under punishment, $\underline{J}(i^*)(1 - \beta)$ and $\underline{U}(i^*)(1 - \beta)$, remain constant, and*
2. *The principal's and the agent's flow payoffs under cooperation, $\overline{J}(i^*)(1 - \beta)$ and $\overline{U}(i^*)(1 - \beta)$, increase and approach $-\pi_a(\eta)\chi$ and $-\eta$, respectively.*

Proposition 6 follows from the fact for a given difference in flow payoff between cooperation and punishment (i.e., holding $(\overline{U}(i^*) - \underline{U}(i^*))(1 - \beta)$ fixed), an increase in β increases the left hand side of (13). Hence, it is possible to provide incentives while forgiving the agent more often (i.e., increasing $\tilde{s}(i^*)$). Since intervention takes place with lower likelihood, this increases the flow payoff from cooperation for both the principal and the agent. In the limit, $\tilde{s}(i^*)$ approaches $\bar{s}$ so that intervention almost never happens and the principal and the agent receive a flow payoff from cooperation equal to the value associated with permanent high effort by the agent.⁴⁷

Note that while the flow payoff under cooperation increases with the discount factor, the same is not true of the flow payoff under punishment. This is because the principal's flow payoff is fixed at $\underline{J}(1 - \beta)$ from (16), which means that the flow payoff to the agent under punishment cannot change either.⁴⁸

6 Extensions

Note that our simple benchmark model ignores two additional issues. First, as we mentioned, it ignores the possibility that the principal can reward the agent for reducing disturbances using either temporary transfers or permanent concessions. Second, it ignores the possibility that the principal can replace the agent. These issues are discussed in the below two extensions which show that our main conclusions are unchanged.⁴⁹

⁴⁶In other words, if the level of intensity exceeds i^* , then it is not possible to decrease intensity while increasing duration, since duration is already at the maximum.

⁴⁷This relies on $\tilde{s}(i^*)$ approaching $\bar{s}$ sufficiently quickly, which is guaranteed by the asymptotic informativeness of s_t implied by $\lim_{s_t \rightarrow \bar{s}} \Phi_s(s_t, 0) / \Phi_s(s_t, \eta) = \infty$.

⁴⁸It is nonetheless true that the duration of punishment $\tilde{d}(i^*)$ approaches 1, though this effect is counterbalanced by an increase in the continuation value under cooperation, and this keeps $\underline{U}(i^*)(1 - \beta)$ unaffected by β .

⁴⁹Due to space restrictions, we describe these results informally, but more details are available upon request.

6.1 Temporary Transfers and Permanent Concessions

6.1.1 Temporary Transfers

Our benchmark model ignores the presence of transfers from the principal to the agent which are standard in principal-agent relationships. A natural question concerns when a government should use transfers and when a government should use interventions in providing the right incentives for the agent.

Consider an extension of our model where if the principal does not intervene at t ($f_t = 0$), he chooses a transfer $\tau_t \geq 0$ 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 $-\tau_t - s_t\chi$ and the payoff to the agent is $\tau_t - e_t$. As in the benchmark model, the value of τ_t chosen by the principal can depend on the entire public history of the game.

We find that the efficient sequential equilibrium of this extended model has the following properties. First, the prospect of future transfers serves as a reward for the successful avoidance of large disturbances.⁵⁰ Second, the use of intervention continues to serve as a punishment for large disturbances. Moreover, transfers are never used during intervention since the principal would like to make the agent suffer as much as possible. Therefore, if a sufficient number of large disturbances occur along the equilibrium path, then intervention occurs exactly as in the benchmark model, and all of our results regarding the dynamics of intervention are completely unchanged.

The main difference between the benchmark model and the extended model is that under some conditions, the first part of Proposition 2 does not hold. In this case, the extended model admits a second long run equilibrium in which intervention is not used.⁵¹ In this long run equilibrium, the principal does not use intervention, and he only uses transfers to provide incentives. More specifically, the long run equilibrium features a transfer phase in which the principal pays the agent and a no-transfer phase in which the principal does not pay the agent. In both phases, the principal requests high effort from the agent. The realization of a small enough disturbance leads to a probabilistic exit from the no-transfer phase and the realization of a large enough disturbance leads to a probabilistic exit from the transfer phase.

Thus, the equilibrium of the extended model can feature history-dependence in the long run contract. On the one hand, the absence of large disturbances along the equilibrium path can lead to an equilibrium which features no intervention and repeated transfers.⁵² On the other hand, too many large disturbances along the equilibrium path can lead to an equilibrium which features no transfers and repeated intervention as in our benchmark model.

⁵⁰This is of course assuming that transferring resources is cheaper for the principal relative to allowing low effort by the agent.

⁵¹This requires the discount factor to be sufficiently high so as to guarantee the existence of a trigger-strategy equilibrium in which transfers induces high effort. Absent this condition, the unique long run equilibrium is the one we have previously described with repeated intervention.

⁵²This is also the case if the initial condition U_0 is chosen to be sufficiently high.

As an aside, note that the equilibrium realization of interventions in this environment requires the existence of a lower bound on the transfer to the agent τ_t . To see why, suppose that transfers were unbounded from below (i.e., the agent can transfer an arbitrarily large amount of resources to the principal). In this situation, the principal could induce permanent high effort by the agent without any intervention by requesting higher and higher payments from the agent as a punishment for the realization of large disturbances. Importantly, these payments could become arbitrarily high following some sequence of disturbances. The use of transfers here would always dominate intervention as a form of punishment since intervention is costly for the principal whereas payment by the agent is beneficial to the principal. This highlights the fact that interventions in our environment occur precisely because of the lower bound on the transfer that the principal can make to the agent.

6.1.2 Permanent Concessions

A government may also provide incentives for the agent in the form of a permanent concession. A natural question concerns how and when a government should provide such a concession. 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 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 $J^C \in (\underline{J}, -\pi_a(\eta)\chi/(1-\beta))$.

We find that the efficient sequential equilibrium of this extended model has the following properties. First, the provision of this concession serves as a reward for the successful avoidance of large disturbances.⁵³ Second, the use of intervention continues to serve as a punishment for large disturbances. Therefore, if a sufficient number of large disturbances occur along the equilibrium path, then intervention occurs. Alternatively, if sufficiently small disturbances occur, then the principal makes a concession which ends all conflict so as to reward the agent for good behavior.

⁵³This is because rewarding the agent by allowing low effort is inefficient for the principal as well as the agent.

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 considered. As in the extension in Section 6.1.1, the long run equilibrium depends the initial point U_0 and on the history of disturbances. All of our results regarding the long run equilibrium under intervention are preserved with one difference. The equilibrium is not quantitatively identical to the one in the benchmark model. This is 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. Given our description of the equilibrium in Proposition 3, this implies that the agent's continuation value under punishment $\underline{U}(i^*)$ must be higher in the extended model. Thus, the likelihood of intervention is higher and its duration shorter because it is harder to provide incentives to the principal and to the agent.⁵⁴

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 in this case 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 quantitatively identical as in our benchmark model.

6.2 Replacement

Our model additionally ignores the possibility that the principal can replace the agent with another identical agent via assassination or demotion. A natural question concerns when a government should use intervention and when a government should use replacement as a threat to provide incentives to the agent.

To explore this question, suppose that at the beginning of every period, the principal can replace the agent, where replacement provides the agent with a continuation value U^R , where for simplicity we assume that U^R is strictly below $\underline{U}$ in the equilibrium which does not allow for replacement. Replacement entails an exogenous cost $\phi \geq 0$ borne by the principal, capturing the cost of removal of the incumbent or training of a replacement agent.⁵⁵ Our benchmark model is embedded in this extended model for ϕ sufficiently high. In that situation replacement is very costly to the principal, and it is never chosen along the equilibrium path since it is strictly dominated by direct control. Moreover, it is clear that if $\phi = 0$, then intervention is never used as a form of punishment since it is strictly dominated by costless replacement. In this situation, our extended model is analogous to the classical Ferejohn (1986) model of electoral control, with the exception that we consider history-dependent strategies. More generally, one can show that

⁵⁴This result is obtained under an analogous condition to Assumption 3, so that the implied likelihood and duration of intervention are interior.

⁵⁵In 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. Specifically, any out of equilibrium removal of an incumbent can prompt all future agents to punish the principal by exerting zero effort forever.

there is a cutoff for the cost of replacement ϕ^* below which replacement serves as the unique form of punishment and above which intervention is the unique form of punishment. Thus, our model coincides exactly to the case for which the cost ϕ exceeds the cutoff.

7 Discussion

7.1 Application: Counterinsurgency

As discussed in the introduction, there are many applications of our model. A particularly relevant application to current affairs is counterinsurgency policy.⁵⁶ 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.⁵⁷ This principle is first outlined in Galula (1963). In this seminal work he suggests that setting up indirect control relationships might be helpful:

"[The counterinsurgent] may, at the same time, utilize to the utmost those who are willing to support him actively, giving them increased privileges and power, and ruling through them, however disliked they may be." (p.102)

Similarly, he suggests that a counterinsurgent can obtain the collaboration of the local leadership with the implicit threat of military intervention:

"The general line could be: stay neutral and peace will soon return to the area. Help the insurgent, and we will be obliged to carry on more military operations and thus inflict more destruction." (p.109)

Our model of indirect control is thus relevant for counterinsurgency policy. Specifically, the use of military interventions in this scenario is an important issue in policy discussions. Indeed, some experts have defended the use of punitive interventions. For example, military strategist Luttwak (2007) writes:

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

⁵⁶In this application, disturbances correspond to attacks by insurgents. See footnote 14 for how one can model the incentives of the insurgents in our framework.

⁵⁷See Nagl (2002) for a discussion.

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 authority—are optimal. More specifically, it shows that this requires the presence of two frictions: double-sided lack of commitment and asymmetric information. It also requires certain additional assumptions. For example, it is necessary that the local authority 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 on the margin to the local authority ($g'(0)$ is sufficiently high) since otherwise temporary costly intervention is suboptimal. Finally, our extensions of Section 6.1.1 and 6.1.2 suggest that even if temporary and costly interventions are sometimes optimal, they need only be used *if a sufficient number of large disturbances have occurred*. Otherwise, the optimal policy is to provide incentives in the form of rewards, either in the form of transfers or in the form of 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 suboptimal, 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 5.2 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. Therefore automatic knee-jerk reactions after every disturbance are a signal of suboptimal conduct. In addition, the analysis of Section 5.3 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 5 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 5 also predicts that an optimally behaving government will also respond with a higher frequency of intervention (higher $l(i^*)$). In sum, international pressure alone cannot remove the need for intervention, and it can have the unintended consequence of making them more frequent. Nonetheless, to the extent that the international community can play a role, the extension in Section 6.1.2 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).

7.2 Historical Examples

This section briefly discusses three historical examples of indirect control. The purpose of this section is to show that the basic elements of indirect control situations can be identified in real-life examples, and that behavior qualitatively conforms to the main predictions of the efficient equilibrium of the model. In the following examples we abstract from many details of these historical episodes in order to bring to the fore the broad patterns of indirect control and punishment expeditions. We discuss common themes across examples and some of the limitations of our model in the light of the examples in the closing subsection.

7.2.1 Early Imperial Rome

Methods of indirect control have been historically often used by imperial powers. Here we consider the historical example of early imperial Rome, from Augustus to Nero, approximately from 20BC to 60AD to exemplify the origin, prevalence, uses, and characteristics of such arrangements. Afterwards, we look into more detail at the German frontier during this period, where the cycle of indirect control, transgressions, and punishments is described through the lens of our model.

During this period the Empire consolidated the vast conquests made by Julius Caesar and the late Republic. The most startling fact about the security system set up is that control over this vast territory was ensured with a very limited amount of troops.⁵⁸ The land borders of this empire were very long: they followed the Rhine in northern Europe, cut across along the Danubius to the Black Sea, and continued south in the Levant along not well-defined geographical landmarks. The same lack of easily defensible positions was true in Northern Africa. These tens of thousands of miles could not be effectively defended by a static deployment of troops, and indeed they were not. This task was outsourced to a vast array of *client states* and *client tribes* that were in a situation of indirect control with respect to the imperial power.⁵⁹

Two features of these client arrangements are noteworthy. First, they were extremely prevalent: From Mauretania in Africa, to Thrace in the Balkans, to dozens of small kingdoms in Anatolia and the Levant (from Armenia to Judea and Nabatean Arabia). Even in the Italian peninsula, two kingdoms ruled over the Celtic inhabitants of the valleys surrounding important trade routes (Alpes Cottiae and Alpes Maritimae). Second, these arrangements straddled a wide spectrum of relationships between the Roman overlords and the local rulers and peoples. At one extreme, local kings were considered little more than appointed officials. At the other extreme, the kingdoms or tribes kept most of their sovereignty and had to be actively courted, incentivized or coerced. This was particularly true of tribal clients across the Rhine who ultimately had the

⁵⁸During most of the period the number of active legions was kept below thirty. A legion consisted of about 6,000 men. In addition, there were *auxilia* formed by non-citizens that Tacitus puts at roughly the same number (*Annals*, IV). This adds up to no more than 350,000 troops. See also Syme (1933) and the modern treatment of Keppie (1996) for an overview of the army and navy of the Roman Empire during this period.

⁵⁹See Luttwak (1976) for the original description of this strategic position. Sidebottom (2007) provides an extended analysis of clients and their relationship with the central Roman power.

option of moving away to avoid Roman interference.⁶⁰ In fact, some of these tribal clients were formed by former or current enemies that were kept in line with a combination of payments and threats of force.⁶¹

The main task that the Romans wanted the clients to perform was to keep internal and external security. Good clients would perform these tasks very efficiently and this way save Roman manpower that could be deployed elsewhere. A good example is King Herod of Judea. During his reign no Roman troop was deployed in an area that would become later on troublesome and need the presence of up to three legions. These savings on the direct use of Roman troops were what ultimately allowed control over such large tracts of land, and protection of long borders. Indeed, the Roman Legions were offensive strike forces that were not designed for defensive warfare. Rather, they were the ultimate coercive tool that the Emperor could deploy in order to force clients to perform their security duties. The fact that the legions were mostly used as a latent threat, meant that a single legion could simultaneously give incentives to many rulers.⁶²

How often the threat of force needed to be realized depended on the internal structure of the clients. In the East, dynastic rulers were very familiar with the power that Rome could deploy, and tended to be well-organized internally. As a consequence, interventions were only necessary when dynastic rivalries or succession struggles threatened stability.⁶³ Also, in this case, rewards to good rulers tended to be of a personal nature—territorial or monetary—and punishments also typically took the form of personal demotion. In contrast, tribal leaders of the West had loose structures of control so the threat of force had to be applied more often and punishment had to be widely applied to tribal populations. Nonetheless, the Romans did their best to keep their relationships with a few noble families, arguably to provide stronger dynamic incentives.⁶⁴ Since the relationship with the Germanic tribes of the West is closer to the focus of the model, the next subsection analyzes it in greater detail.

The Case of Germania

After the conquest of Gallia, Julius Caesar had to face the fact that Germanic tribes had been drifting West across the Rhine and pillaging the inhabitants of the newly conquered Roman provinces for generations. He crossed the Rhine twice (55BC and 53BC) to punish tribes that had looted territories West of the river.⁶⁵ He conquered and stabilized the left bank of the Rhine

⁶⁰This is what the Marcomanni under King Maroboduus did in the last decade BC. See Dobiáš (1960).

⁶¹The examples of the Cherusci and Batavi Germanic tribes appear in Tacitus' *Germania* and the *Annals*.

⁶²"Sometimes dependent and therefore obedient, and sometimes hostile, client tribes and tribal kingdoms required constant management with the full range of Roman diplomatic techniques, from subsidies to punitive warfare." (Luttwak, 1976, p.36). See also Sidebottom (2007), Millar (1993), and Richardson (1991).

⁶³An example are the interventions required in Judea after Herod's succession problems with his sons (see Josephus, 1959).

⁶⁴See, for instance, the example of the family of Arminius in the next subsection.

⁶⁵"The German war being finished, Caesar thought it expedient for him to cross the Rhine, for many reasons; of which this was the most weighty, that, since he saw the Germans were so easily urged to go into Gaul, he desired they should have their fears for their own territories, when they discovered that the army of the Roman people both could and dared pass the Rhine. There was added also, that portion of the cavalry of the Usipetes and the Tenchtheri, which I have above related to have crossed the Meuse for the purpose of plundering and procuring forage, and was not present at the engagement, had betaken themselves, after the retreat of their countrymen,

and the frontier was afterwards maintained with this system of punitive expeditions.

Under Augustus, this situation changed in 12BC when Drusus crossed the Rhine in force to conquer and pacify the land between the Rhine and the Elbe. After a series of battles he decisively weakened Germanic resistance. He was succeeded in this effort by his brother (and future Emperor) Tiberius. By 6BC the area between the two rivers was considered under control, although restless, and several relationships were established with the tribes present.⁶⁶ The Roman presence, in any case, was very scarce, so the relationship is better understood as one of indirect control.

This situation of indirect control was punctuated by the "Varian disaster." In 9AD three full legions under the command of Quinctilius Varus were ambushed and destroyed in a large uprising of the Cherusci commanded by Arminius. Arminius was a noble of the tribe that had been elevated to Roman citizenship. He thus was effectively supposed to be an agent of Roman control who clearly failed to exert effort in keeping the peace.

In the model, this uprising is an example of a very high s_t realization, which requires punishment. The punishment took the form of an extremely brutal two-year campaign in 14-16AD commanded by Germanicus.⁶⁷ Interestingly, in 16AD emperor Tiberius recalled Germanicus from the front thus halting the punishment expedition. He explicitly stated that he considered diplomacy better than war in obtaining cooperation:

He himself had been sent nine times to Germania by the deified Augustus; and he had effected more by policy than by force. Policy had procured the Sugambrian surrender; policy had bound the Suebi and King Maroboduus to keep the peace. (Annals II.26).

With this decision the situation of indirect control was restored with the Empire's frontier on the Rhine, and with the understanding with several tribes across the river that they would be held responsible if raids were launched into Roman territory from their lands. This conforms, in the optimal equilibrium of the model, to a return to the long term phase of cooperation after a long phase of punishment. These cycles, somewhat attenuated, continued. Cross-Rhine Germanic raids followed by Roman punitive expeditions are documented in 21AD, 42AD and 50AD.⁶⁸ However, the underlying continuity of indirect control is exemplified by the fact that a nephew of Arminius himself was elevated by the Romans as King of the Cherusci.⁶⁹

across the Rhine into the territories of the Sigambri, and united themselves to them." (De Bello Gallico IV.16) This paragraph shows that the purpose of these expeditions was to punish past transgressions and to show that further future punishment was possible.

⁶⁶"The Gallo-German nobility on both sides of the Rhine, whose allegiance Rome as the occupying power sought to secure through individual grants of citizenship and absorption into the ranks of *equites*, was a pillar of Romanization." (Rüger, 1996, p.528)

⁶⁷See Shotter (1992).

⁶⁸See Rüger (1996) and Annals XIII.56.

⁶⁹See Luttwak (1976).

7.2.2 Israel in the Palestinian Territories

In this section, we consider the historical example of Israeli policy in the Palestinian Territories following the Oslo Accords of 1993. This set of agreements put Israel and the Palestinian Authority (PA) in an explicit relationship of indirect control.⁷⁰ More specifically, under this arrangement, Israel would free areas from military occupation in exchange for the PA's agreement to exert the highest effort in minimizing terrorist attacks against Israel from these areas.⁷¹

As predicted by the model, the PA was expected to exert the needed effort, but asymmetric information prevented full trust. While the extent to which the PA consistently exerted effort is obviously unknown, there are many instances in which visible actions were taken. For example, 1,200 suspected Islamists were arrested, the Islamic University and some thirty Hamas institutions were raided, and the Gaza mosques were put under PA control following a string of suicide bombings in Tel Aviv and Jerusalem in 1996. There are other examples of such crackdowns, and also rumors that the PA cooperated with the Israeli Defense Forces by providing information on the location of Hamas and Islamic Jihad activists throughout the 1990s.⁷² These crackdowns were at times important enough to create a rift within Palestinian society, and to receive praise from Israeli and US officials.⁷³ Nevertheless, the actual extent of PA cooperation was unclear throughout the period.⁷⁴

Up to 1996, Israel followed a policy of carrot and stick to encourage PA effort. In response to terrorist attacks, further steps in the development of the Oslo accords would be frozen.⁷⁵

⁷⁰Jamal (2005) writes: "This policy of strict control over all realms of life continued until the establishment of the PA in 1994; then the occupied territories were divided into three areas with different legal status, and Israeli control of the West Bank and Gaza Strip was transformed from direct to indirect" (p. 29). See also Kristianasen (1999) and Said (2000).

⁷¹Beinin (2006) writes: "Rabin initially saw the Declaration of Principles as a security arrangement. Shortly before its approval he explained:

I prefer the Palestinians to cope with the problem of enforcing order in Gaza. The Palestinians will be better at it than we because they will allow no appeals to the Supreme Court and will prevent the Association for Civil Rights from criticizing conditions there by denying it access to the area." (p. 29)

See also Said (1995).

In the interim agreement on the West Bank and Gaza reached in 1995 (known either as Oslo II or Taba Accords) it is explicitly stated: "Except for the Palestinian Police and the Israeli military forces, no other armed forces shall be established or operate in the West Bank and the Gaza Strip." The PA was thus charged with uprooting armed factions. These security guarantees were even more explicit in the Wye River Memorandum of 1998, where the PA was again required to outlaw and combat terrorist organizations.

⁷²See Kristianasen (1999).

⁷³In a New York Times article, Ibrahim (1994) reports "In Gaza and increasingly in the West Bank, Palestinians who once regarded Israel as the sole enemy have come to see the Palestine Liberation Organization and its chairman, Mr. Arafat, as another enemy." Also, in a later article, Sciolino (1995), reports "Mr. Christopher—Secretary of State of the United States—said Mr. Arafat had made a "100 percent" commitment to bring terrorists to justice." Finally, Greenberg (1996) reports Peres saying: "No doubt the Palestinian Authority has prevented a few cases of infiltration into Israel."

⁷⁴Newspapers reported rumors that "After a terrorist attack against Israel, the Palestinian police arrest Islamic fundamentalists, who are quietly released a few weeks later" (Halevi, 1996).

⁷⁵For instance, Greenberg (1996) reports "Mr. Peres had suspended contacts with the Palestinians and delayed an army withdrawal from most of Hebron after the suicide attacks, demanding that Mr. Arafat crack down on

However, progress in agreement implementation would resume and further concessions would be granted following the absence of attacks.⁷⁶ In the context of our model, this is consistent with the transitory period in the optimal equilibrium where incentives are provided without actual use of punishment.

This two-pronged approach came to a halt in mid-1996, with the arrival in power of Netanyahu, who defeated Peres by pointing out that the peace process had not stopped terrorist attacks.⁷⁷ While our model cannot possibly capture the intricacies of the electoral contest in Israel, the outcome of the election and the subsequent absence of further concessions is consistent with the path of play reaching the long term equilibrium in our model.⁷⁸ From 1996 onwards there was a steady increase in Israeli military intensity and punitive measures, such as house demolitions and assassinations. This rise culminated in the restoration of military control over the entirety of the West Bank in 2002, as a belated response to the explosion of the second intifada and the subsequent increase in terrorist attacks, which can be seen as a large realization of s_t .⁷⁹ The punishment dimension of this intervention is sometimes openly discussed.⁸⁰

Our comparative statics from section 5.3 suggest that the model may guide us in understanding this steady increase in intensity.⁸¹ More specifically, there are three parameter changes which can result in such outcome. First, and most obvious, the model predicts that an increase in intensity follows an increase in χ , the cost to Israel of a Palestinian attack. The increasing use of suicide bombings by Hamas and Islamic Jihad throughout the 1990s might thus explain the rise in the Israeli use of force. Moreover, following Ariel Sharon's visit to the Temple Mount in September 28, 2000, there was a dramatic increase in the number of terrorist attacks as part of the al-Aqsa intifada.⁸² Such increase in the deadliness and frequency of terrorist attacks is

Islamic militants."

⁷⁶For instance, Peres took steps to limit the expansion of settlements in the West Bank and organised the first Presidential and legislative elections in Palestinian history in the West Bank. See Quray (2008, p. 12). Also, after the Palestine National Council amended clauses of the Charter of the PLO that called for the destruction of Israel, the election platform approved by the Labour Party did not rule out anymore a Palestinian state.

⁷⁷As Quray (2008, p. 12) describes it, Peres lost the election after "a series of suicide bombings by Palestinian Islamist groups in February and March 1996 fatally undermined his authority."

⁷⁸While dialogue and bargaining officially continued, few further concessions were actually implemented. According to Quray (2008), when Netanyahu left office in 1999, "the peace process [was] almost extinct."

⁷⁹According to Enderlin (2006, p.8) there was a change in approach: "Individual soldiers allowed themselves to react more spontaneously. They no longer feared being the target of an inquiry by the military policy each time a Palestinian civilian was killed. Beginning in November 2000, the IDF officially designated the intifada an "armed conflict: combat against terrorist groups". New procedures were put in place. The military police no longer immediately investigated the circumstances of a civilian death." According to this source, of the 3,185 civilian deaths which occurred between September 2000 and June 2005, only 131 were investigated and 18 resulted in indictments.

⁸⁰In Enderlin (2006, p. 36) Gal Hirsh-IDF military-is quoted saying that "The operations of the Israeli army aimed to demonstrate to the Palestinian Authority that it would pay the price for its support of terrorism." Also, in p. 12 Lieutenant General Yaalon says "It is of the utmost importance that the war conclude on the affirmation of a principle, which is that the Palestinians realize that violence does not pay."

⁸¹This is the case subject to the caveat that our model only allows us to compare across steady states and does not shed light on the transition path from one steady state to another.

⁸²See Hammami and Tamari (2006).

therefore in line with the rise in Israel's intensity of intervention.⁸³

Second, the model predicts that an increase in η , the cost to the agent of limiting disturbances, is also associated with an increase in the intensity of intervention. This cost can increase due to a loss of legitimacy of the agent, or due to an increased preference for attacks by the agent (or the population he is representing). These two forces were present in the Palestinian territories. The perception that Israel was not keeping up its side of the bargain, mostly due to the growth in settler population, together with the rampant corruption in the PA administration both increased the popularity enjoyed by Hamas, and with it the support for terrorist activities. In December 1995, 77.9% of Palestinians supported the peace process, but such support steadily declined and was only 44.1% in December 1999.⁸⁴

Finally, the model also predicts an increase in intensity if there is a reduction in the marginal cost of violence, A . After 9/11 international public opinion and in particular American opinion became more tolerant of heavy handed action against terrorism.⁸⁵ To the extent that international rebuke is a large component of A , such changes in attitudes may have contributed to the rise in military intensity by Israel. In sum, the increase in intensity in the cycles of military intervention from 1996 onward is consistent with the changes observed in three exogenous variables in the light of the model.⁸⁶

7.2.3 Chechen Wars

Indirect control can also be observed in the current relationships between actors within states, which often feature bouts of recurring violence, as in the case of Chechnya.

In the aftermath of the dissolution of the USSR, Russia faced the need to establish new relationships with all the territorial elements within its borders. This was eventually achieved with all regions except Chechnya. This territory unilaterally declared independence in late 1990, and by 1991 was organizing under the leadership of Dzhokhar Dudaev, who emerged victorious in a series of internal struggles. While the Russian state did not accept this new situation, it did not immediately intervene. Thanks to this ambiguous situation, Chechnya became a "free economic zone" where smugglers, arms traders, oil thieves, hijackers and tax evaders found a haven and markets.⁸⁷ In addition to these negative externalities, Russian authorities were concerned that this de facto independence could influence other regions.⁸⁸ To try to solve these

⁸³See Baliga and Sjöström (2009) for an interesting model of provocateurs that incite escalation.

⁸⁴Data from the Jerusalem Media and Communication Center, as cited in Jamal (2005, p. 151). On the steady erosion of PA popularity leading to the outbreak of the second intifada, see also Hammami and Tamari (2001).

⁸⁵When asked in a Time/CNN survey days after the attacks, 41% reported feeling less favorable toward Palestinians as a result of 9/11, and just 3% felt more favorable. This information is available at http://www.americans-world.org/digest/regional_issues/IsraelPalestinians/viewIsrPal.cfm

⁸⁶As a caveat, we cannot claim that Israel's use of military intervention was itself optimal or that its intensity was optimally chosen. To make such statements one would have to argue that the conditions outlined in the previous subsection (including whether a sufficient number of large disturbances occurred before intervention, and whether the use of positive incentives such as territorial concessions was contemplated and used) were satisfied.

⁸⁷For a description of the situation and its negative consequences, see Gall and de Waal (1997).

⁸⁸See, for instance, Evangelista (2002), p. 86.

problems Russia tried to replace Dudaev with a more amenable proxy by giving progressively more direct support to the Chechen opposition, who was responsible for at least three attempts to militarily overthrow Dudaev.⁸⁹ These operations failed and as a consequence in 1994 the Russian army invaded Chechnya.

This invasion developed into a full-fledged war that lasted until 1996, at great cost to both the Chechen population and the Russian military. About 50,000 Chechens are estimated to have died during this conflict, mostly as a consequence of aerial bombing that reduced the main cities to rubble.⁹⁰ While some in the Russian military were prepared to continue the intervention, in the Khasavyurt agreements of 1996 Russia accepted the local authority of Aslan Maskhadov, Dudaev successor, and withdrew the army.⁹¹ In these accords the decision over the final status of Chechnya was deferred for 5 years during which Chechnya was supposed to remain part of a “common economic space.”⁹²

This new status failed to satisfy Chechen radicals such as Basaev, Raduev, and other Islamists, who started a destabilizing campaign of violence both inside Chechnya and in Russia proper. For example, in 1997 there were bomb attacks in Armavir and Nalchik (ordered by Raduev) and two British citizens were kidnapped to boycott oil agreements between BP, Russia, and Chechnya.⁹³ However, Russia did not intervene directly to answer these attacks. Instead it put pressure on Maskhadov who was thus forced to confront the radicals.⁹⁴ Hence, in the light of our model, Russia used Maskhadov’s effort as an agent to indirectly control Chechen terrorists for about three years.⁹⁵

This indirect control situation came to a head in August 1999, when the Islamic militias of Basaev and Khattab invaded neighboring Dagestan from their bases in Chechnya, at the same time that terrorist attacks took place in Moscow and several other cities. This can be interpreted as a very large realization of s_t , which demands punishment. The Russian army defeated the insurgents in Dagestan and proceeded to attack Chechnya proper while Putin gave an ultimatum to Maskhadov “to arrest those responsible for the invasion of Dagestan or face further attacks.”⁹⁶ Note that this ultimatum makes sense in the context of implicit indirect control that had prevailed during the interwar years.

The attacks on Chechnya escalated into the second Chechen war. This war caused again very heavy civilian casualties and destruction, to the point that some accused the Kremlin of “waging

⁸⁹See Hughes (2001), p. 30.

⁹⁰In addition, there are extensive reports of widespread use of torture and arbitrary detentions and executions in filtration camps (See for instance Gall and de Waal, 1997, Evangelista, 2001 and Wood, 2007)

⁹¹For a sign that some in the military were willing to continue the fight, note that General Pulikovsky was announcing a renewed offensive over Grozny only days before the signing of the accords (See Wood, 2007, p. 75).

⁹²See Hughes (2001), p. 32 for a description of the accords.

⁹³See Evangelista (2002), p. 52.

⁹⁴For example, several days of fighting took place in 1998 between government forces and Islamic paramilitary units in Gudermes (see Wood, 2007, p. 91). Also, a warrant to detain Raduev was issued in 1997 after he claimed credit for bomb attacks in Russia (Evangelista, 2001, p. 52).

⁹⁵Interestingly, some Russian politicians such as Lebed, signatory of the Khasavyurt Accords, criticize Yeltsin on the basis that he was not giving enough support to Maskhadov (See Evangelista, 2001, p. 57).

⁹⁶See Evangelista (2002), p. 68.

a war against the Chechen people” (Trenin and Malashenko, 2004, p. 41).⁹⁷ The objective of this new war was quite obviously to install a new pro-Russian proxy ruler, Akhmad Kadyrov, solidly in power. Kadyrov was appointed head of the Chechen administration in 2000 and had expressed strong willingness to fight the Islamist rebels.⁹⁸

As Russian troops have gradually withdrawn from the region, the Kremlin has followed a strategy of “Chechenization” which puts Chechnya again in a situation of indirect control with Ramzan Kadyrov, the son of the assassinated Akhmad, in charge of controlling pro-independence and Islamist forces.⁹⁹

7.2.4 Summary and Further Extensions

In these three examples, we have discussed how the descriptive patterns in these conflicts can broadly conform to the setup and the characterized equilibrium of the model. These three examples share three important features related to the setup and results of our model.

First, each situation corresponds to one of indirect control with some asymmetric information. More specifically, in all situations, agents have the capacity to facilitate or hinder disturbances which cause discomfort to the principal. These agents are implicitly tasked to do this by the principal and are held responsible if these disturbances occur. Along the equilibrium path, insofar as the observer can tell, agents are exerting effort in thwarting disturbances. However, uncertainty remains regarding whether this effort is enough or is appropriate.

Second, each situation features the occasional use of military interventions by the principal. More specifically, the principal is significantly stronger militarily than the agent, and he utilizes this asymmetry in order to provide incentives to the agent. Along the equilibrium path, the principal forgives many transgressions during the transition to the cycle of repeated intervention. Also, during this transition, the absence of disturbances can be rewarded with concessions. However, military intervention eventually occurs after large disturbances (high s_t realizations).

Finally, all of the interventions are temporary, repeated, and are often deemed excessively destructive by outside observers. After a period of punishment, the principal withdraws and re-establishes the situation of indirect control. Our model explains these patterns by showing that it is precisely because the interventions are excessive (i.e., statically inefficient) that they are temporary and repeated, since this is the only credible means for the principal of providing incentives to the agent.

This brief discussion of historical examples also highlights some of the limitations of the model. One limitation is that the model analyzes a stationary environment and finds that indirect control is optimal. Hence the model can not explain how a situation of indirect control is established or how this situation ends. Factors outside the model thus must be used to explain, for instance, why the Oslo accords are signed in 1993 and not before, or after. Similarly, wars of

⁹⁷ There were 25,000 civilian casualties according to Amnesty International, 2007.

⁹⁸ See Politkovskaya (1999), p. 197.

⁹⁹ See Wood (2007) for a description of the current situation.

decolonization or the war of Eritrean independence are examples of abandonment of a situation of indirect control that the model currently has difficulty fitting.¹⁰⁰

A second related limitation is that the model characterizes a steady state with repeated intervention in which the intensity, likelihood, and duration of intervention are fixed constants. Obviously, in reality these change over time, since the intensity of force can rise or decline through the course of an intervention or at different points in the cycle. This occurs in part because many of the exogenous features of our environment—such as the level of effort required by the agent to control disturbances—are also changing. A richer model would incorporate these time varying features in the environment.

Third, as the Chechen case makes clear, not any agent seems equally amenable to a relationship of indirect control. For this reason, principals often actively try to replace and install particular agents as proxy rulers, and equilibrium replacement and interventions often take place together in ways that differ from our extension in Section 6.2. The main reason these phenomena do not occur in our model is because all agents are identical.

The limiting simple structure of our predictions is driven in large part by the fact that, in our model, hidden information is one-sided and i.i.d. This is done in order to preserve the simple repeated game structure of the model. In practice, hidden effort by the agent in preventing disturbances can often have a long-lasting effects, and both the agent and the principal can have persistent hidden information about their type. For instance, the cost of high effort by the agent can be hidden information as well as the cost of the principal’s intensity of intervention. One can conjecture that in an environment which incorporates these features, interventions would not necessarily take the simplistic two phase form that we have described, and the intensity, likelihood, and duration of intervention would vary over time. Moreover, persistent hidden information on the side of the agent would lead to an eventual motive for the principal to replace an agent who is deemed an undesirable type, implying that both equilibrium interventions and replacement would coexist. Finally, the process of learning about principal and agent types could potentially also explain when and why indirect control situations are established and abandoned. Hence, we hypothesize that adding some dimensions of persistent asymmetric information to the model can generate more complex dynamics and expand further the explanatory power of the principal-agent framework vis à vis relationships of indirect control.

8 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 generates recurring bouts of violence and highlights a fundamental tradeoff between the intensity and duration of interventions. It also allows us to consider the separate effects of a rise in the

¹⁰⁰Clearly, the pattern does not always conform to the extension in Section 6.1.2 where permanent concessions are offered in return for an extended period of good behavior.

cost of effort to the agent, a fall in the cost of intensity to the principal, and a rise in the cost of disturbances to the principal.

We have also discussed the setting of the model and its main predictions in the context of three historical examples. Our model sheds new light on the forces behind the cycles of violence and on other aspects of indirect control relationships. However, it also abstracts from a number of potentially important issues, as our discussion of historical examples makes clear.

Our research thus suggests several avenues for further research. 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, as noted in the discussion of the historical examples, 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.

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

Acemoglu, Daron and Alexander Wolitzky (2011) "The Economics of Labor Coercion," *Econometrica* 79, 555-600.

Ambrus, Attila and Georgy Egorov (2009) "Delegation and Non-monetary Incentives", Working Paper.

Amnesty International (2007) "Russian Federation: What Justice for Chechnya's disappeared," Amnesty International Publications 2007.

Anderlini, Luca, Dino Gerardi, and Roger Lagunoff (2009) "Social Memory, Evidence, and Conflict," *Review of Economic Dynamics*, 13, 559-574.

Athey, Susan, Kyle Bagwell, and Chris Sanchirico (2004) "Collusion and Price Rigidity," *Review of Economic Studies*, 71, 317-349.

Baliga, Sandeep and Jeffrey C. Ely (2010) "Torture," Working Paper.

Baliga, Sandeep and Tomas Sjöström (2009) "The Strategy of Manipulating Conflict," Working Paper.

Beinin, Joel (2006) "The Oslo Process and the Limits of a Pax Americana," in *The Struggle for Sovereignty: Palestine and Israel 1993-2005*, edited by Joel Beinin and Rebecca L. Stein., Stanford University Press, Stanford.

Caesar, Julius (2010) *De Bello Gallico and other commentaries*. Bibliobazaar.

Chassang, Sylvain and Gerard Padró i Miquel (2009) "Economic Shocks and Civil War," *Quarterly Journal of Political Science*, 4, 211-228.

Chassang, Sylvain and Gerard Padró i Miquel (2010) "Conflict and Deterrence under Strategic Risk," *Quarterly Journal of Economics*, 125, 1821-1858.

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.

- DeMarzo, Peter M. and Michael J. Fishman (2007)** "Agency and Optimal Investment Dynamics," *Review of Financial Studies*, 20, 151-188.
- Dobiáš, Josef (1960)** "King Maroboduus as a Politician," *Klio*, 38, 155-166.
- Egorov, Georgy and Konstantin Sonin (2009)** "Dictators and their Viziers: Endogenizing the Loyalty-Competence Tradeoff," forthcoming in *Journal of the European Economic Association*.
- Enderlin, Charles (2006)** *The Lost Years: Radical Islam, Intifada, and Wars in the Middle East 2001-2006*, Alpha Graphics, New Hampshire.
- Evangelista, Matthew (2002)** *Will Russia Go the Way of the Soviet Union?*, Brookings Institution, Washington, D.C..
- Farrell, Joseph and Eric Maskin (1989)** "Renegotiation in Repeated Games" *Games and Economic Behavior*, 4, 327-360.
- Fearon, James D. (2004)** "Why Do Some Civil Wars Last So Much Longer than Others", *Journal of Peace Research*, 41, 275-301.
- Ferejohn, John (1986)** "Incumbent Performance and Electoral Control," *Public Choice*, 50, 5-25.
- Fong, Yuk-fai and Jin Li (2010)** "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.
- Gall, Carlotta and Thomas de Waal (1997)** *Chechnya. A Small Victorious War*, Macmillan Publishers Ltd, London.
- Green, Edward J. and Robert H. Porter (1984)** "Noncooperative Collusion Under Imperfect Price Information," *Econometrica*, 52, 87-100.
- Greenberg, Joel (1996)** "Peres and Arafat, After Talks, Agree to Revive Peace Efforts," *The New York Times*, 18 April, Section A, p. 13.
- Guriev, Sergei (2004)** "Red Tape and Corruption," *Journal of Development Economics*, 73, 489-504.
- Halac, Marina (2010)** "Relational Contracts and the Value of the Relationship," Working Paper.
- Halevi, Yossi Klein (1996)** "How Arafat Uses Hamas," *The New York Times*, 14 March, Section A, p.23.
- Hammami, Rema and Salim Tamari (2001)** "The Second Uprising: End or New Beginning?," *Journal of Palestine Studies*, 30, 5-25.
- Hammami, Rema and Salim Tamari (2006)** "Anatomy of Another Rebellion: from Intifada to Interregnum," in *The Struggle for Sovereignty: Palestine and Israel 1993-2005*, edited by Joel Beinin and Rebecca L. Stein, Stanford University Press, Stanford.
- Herbst, Jeffrey (2000)** *States and Power in Africa*, Princeton University Press, Princeton.

- Hughes, James (2001)** "Chechnya: The Causes of a Protracted Post-Soviet Conflict," *Civil Wars*, Vol.4, No.4 (Winter 2001), 11-48.
- Ibrahim, Youssef M. (1994)** "Support for Arafat in Gaza Replaced by Wide Enmity," *The New York Times*, 3 Dec, Section 1, p.1.
- Jackson, Matthew O. and Massimo Morelli (2008)** "Strategic Militarization, Deterrence, and War," Working Paper.
- Jaeger, David A. and M. Daniele Paserman (2008)** "The Cycle of Violence?: An Empirical Analysis of Fatalities in the Israeli-Palestinian Conflict," *American Economic Review*, 98, 1591-1604.
- Jamal, Amal (2005)** *The Palestinian National Movement: Politics of Contention, 1967-2005*, Indiana University Press, Bloomington and Indianapolis.
- Josephus (1959)** *The Jewish War*, translated by G.A. Williamson, Penguin Books, London.
- Keppie, Lawrence (1996)**. "The Army and the Navy" in *Cambridge Ancient History vol 10*. Eds. Alan K. Bowman, Edward Champlin and Andrew Lintott, Cambridge University Press, New York.
- Kristianasen, Wendy (1999)** "Challenge and Counterchallenge: Hamas's Response to Oslo," *Journal of Palestine Studies*, 28, 19-36.
- Levin, Jonathan (2003)** "Relational Incentive Contracts," *American Economic Review*, 93, 835-857.
- 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.
- Mailtah, George J. and Larry Samuelson (2006)** *Repeated Games and Reputation*, Oxford University Press, New York.
- Millar, F. G. B. (1993)** *The Roman Near East, 31 bc–ad 337*. : Harvard University Press, Cambridge.
- 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*, Chicago University Press, Chicago.
- Phelan, Christopher and Robert M. Townsend (1991)** "Computing Multi-Period Information-Constrained Optima," *Review of Economic Studies*, 58, 851-881.
- 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.
- Politkovskaya, Anna (1999)** *A Dirty War*, The Harvill Press, London.

- Powell, Robert (1999)** *In the Shadow of Power: States and Strategies in International Politics*, Princeton University Press; Princeton.
- Quray, Ahmad (2008)** *Beyond Oslo, the struggle for Palestine: inside the Middle East peace process from Rabin's death to Camp David*, I. B. Tauris, London and New York.
- Rabbani, Mouin (2006)** "Palestinian Authority, Israel Rule," in *The Struggle for Sovereignty: Palestine and Israel 1993-2005*, edited by Joel Beinin and Rebecca L. Stein, Stanford University Press, Stanford.
- Radner, Roy (1985)** "Repeated Principal-Agent Games with Discounting," *Econometrica*, 53, 1173-1198.
- Radner, Roy (1986)** "Repeated Partnership Games with Imperfect Monitoring and No Discounting," *Review of Economic Studies*, 53, 43-57.
- Reno, William (1998)** *Warlords Politics and African States*, Lynne Rienner Publishers, London.
- Richardson, J. S. (1991)** "Imperium Romanum: Empire and the Language of Power," *Journal of Roman Studies* 81: 1-9.
- Rogerson, William P. (1985)** "Repeated Moral Hazard," *Econometrica*, 53, 69-76.
- Rohner, Dominic, Mathias Thoenig, Fabrizio Zilibotti (2010)** "War Signals: A Theory of Trade, Trust, and Conflict," Working Paper.
- Rotemberg, Julio and Garth Saloner (1986)** "A Supergame Theoretic Model of Price Wars During Booms," *American Economic Review*, 76, 390-407.
- Rüger, C. (1996)** "The West: Germany," in *Cambridge Ancient History vol 10*. Eds. Alan K. Bowman, Edward Champlin and Andrew Lintott. Cambridge University Press
- Said, Edward W. (1995)** *Peace & its Discontents: Gaza-Jericho 1993-1995*, Vintage, London.
- Said, Edward W. (2000)** *The End of the Peace Process*, Granta Publications, London.
- Sannikov, Yuliy (2007)** "Games with Imperfectly Observable Actions in Continuous Time," *Econometrica*, 75, 1285-1329
- Sannikov, Yuliy (2008)** "A Continuous Time Version of the Principal-Agent Problem," *Review of Economic Studies*, 75, 957-984.
- Sciolino, Elaine (1995)** "Arafat Pledges to Prosecute Terrorists, but Faults Israel," *The New York Times*, 11 March, Section 1,p.5.
- David Shotter (1992)** *Tiberius Caesar*, Routledge, London.
- Spear, Stephen E. and Sanjay Srivastava (1987)** "On Repeated Moral Hazard with Discounting," *Review of Economic Studies*, 54, 599-617.
- Staiger, Robert W. and Frank Wolak (1992)** "Collusive Pricing with Capacity Constraints in the Presence of Demand Uncertainty," *RAND Journal of Economics*, 203-220.
- Syme, Ronald (1933)** "Some Notes on the Legions under Augustus," *Journal of Roman Studies*, 23, 14-33.

Syme, Ronald (1934) "The Northern Frontiers under Augustus," in *Cambridge Ancient History vol 10*, edited by C.T. Seltman, Cambridge University Press, London.

Schwarz, Michael and Konstantin Sonin (2008) "A Theory of Brinksmanship, Conflicts, and Commitments," *Journal of Law, Economics, and Organization*, 24, 163-183.

Sidebotton, Harry (2007) "International Relations," in *Cambridge History of Greek and Roman Warfare vol 2*, edited by Philip Sabin, Hans Van Wees and Michael Whitby, Cambridge University Press, London.

Tacitus, Cornelius (2005) *The Annals of the Imperial Rome*. Digireads.com, Stillwell, KS.

Tacitus, Cornelius (2010) *The Agricola and the Germania*. Penguin Classics, London.

Trenin, Dmitri V. and Aleksei V. Malashenko (2004) "Russia's Restless Frontier. The Chechnya Factor in Post-Soviet Russia," Washington: Carnegie Endowment for International Peace.

Wood, Tony (2007) *Chechnya: the Case for Independence*, Verso, London and New York.

Yared, Pierre (2010) "A Dynamic Theory of War and Peace," *Journal of Economic Theory*, 145, 1921-1950.

10 Appendix A

10.1 Equilibrium Definition

In this section we provide a formal 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 .¹⁰¹ Let $h_t^1 = \{h_t^0, f^t, z^{2t}\}$, 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 expected continuation values for player $j = \{p, a\}$ at h_t^0 and h_t^1 , respectively, as $V_j(\sigma|_{h_t^0})$ and $V_j(\sigma|_{h_t^1})$ 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 $j = p, a$

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

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. With some abuse of notation, define an equilibrium allocation as a function of the exogenous history:

$$\alpha = \{f_t(q_t^0), i_t(q_t^1), e_t(q_t^1)\}_{t=0}^\infty. \quad (19)$$

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 . Moreover, with some abuse of notation, let $V_j(\alpha|_{q_t^0})$ and $V_j(\alpha|_{q_t^1})$ 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.

Lemma 2 $\alpha \in \mathcal{F}$ is a sequential equilibrium allocation if and only if

$$V_p(\alpha|_{q_t^0}) \geq \underline{J} \text{ and } V_a(\alpha|_{q_t^0}) \geq \underline{U} \quad \forall q_t^0, \quad (20)$$

$$V_p(\alpha|_{q_t^1}) \geq \underline{J} \quad \forall q_t^1 \text{ s.t. } f_t(q_t^1) = 1, \text{ and} \quad (21)$$

$$V_a(\alpha|_{q_t^1}) \geq \max \left\{ \begin{array}{l} -\eta + \beta E \left\{ V_a(\alpha|_{q_{t+1}^0}) | q_t^1, e_t = \eta \right\} \\ \beta E \left\{ V_a(\alpha|_{q_{t+1}^0}) | q_t^1, e_t = 0 \right\} \end{array} \right\} \quad \forall q_t^1 \text{ s.t. } f_t(q_t^1) = 0 \quad (22)$$

¹⁰¹Without loss of generality, we let $i_t = 0$ if $f_t = 0$ and $e_t = 0$ if $f_t = 1$. Moreover, we let $s_t = 0$ if $f_t = 1$.

for $\underline{J} = -\pi_p \chi / (1 - \beta)$ and some $\underline{U} \leq -g(0) / (1 - \beta)$.

Proof. The necessity of (20) for $j = p$ follows from the fact that the principal can choose $\hat{f}_{t+k}(q_{t+k}^0) = 1 \ \forall k \geq 0$ and $\forall q_{t+k}^0$ and $\hat{i}_{t+k}(q_{t+k}^1) = 0 \ \forall k \geq 0$ and $\forall q_{t+k}^1$, and this delivers continuation value $\underline{J}$. The necessity of (20) for $j = a$ follows from the fact that the agent can choose $\hat{e}_{t+k}(q_{t+k}^1) = 0 \ \forall k \geq 0$ and $\forall q_{t+k}^1$, and this delivers a continuation value no smaller than some lower bound $\underline{U}$ below $-g(0) / (1 - \beta)$. The necessity of (21) follows from the fact that conditional on $f_t(q_t^0) = 1$, the principal can choose $\hat{f}_{t+k}(q_{t+k}^0) = 1 \ \forall k > 0$ and $\forall q_{t+k}^0$ and $\hat{i}_{t+k}(q_{t+k}^1) = 0 \ \forall k \geq 0$ and $\forall q_{t+k}^1$, and this delivers continuation value $-\pi_p \chi + \beta \underline{J} = \underline{J}$. The necessity of (22) follows from the fact that conditional on $f_t(q_t^0) = 0$, the agent can unobservably choose $\hat{e}_t(q_t^1) \neq e_t(q_t^1)$ and follow the equilibrium strategy $\forall k > 0$ and $\forall q_{t+k}^1$.

For sufficiency, consider a feasible allocation which satisfies (20) – (22) 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. Conditional on q_t^0 , then a deviation by the principal to $\hat{f}_t(q_t^0) \neq f_t(q_t^0)$ is weakly dominated by (20). Moreover, conditional on $f_t(q_t^0) = 1$, a deviation by the principal at q_t^1 to $\hat{i}(q_t^1) \neq i(q_t^1)$ is weakly dominated by (21). If $f_t(q_t^0) = 0$, then a deviation by the agent to $\hat{e}_t(q_t^1) \neq e_t(q_t^1)$ is weakly dominated by (22). ■

10.2 Description of Generalized Problem

Given that we are interested in characterizing (1) – (8) as well as (1) – (8) which ignores (3), (4), and (6), we provide in this section some results which apply to the following generalized problem:

$$J(U) = \max_{\rho} E_z \left\{ f_z \left(-\pi_p \chi - A i_z + \beta J(U_z^F) \right) + (1 - f_z) \left(-\pi_a(e_z) \chi + \beta E_s \left\{ J(U_{z,s}^N) | e_z \right\} \right) \right\} \quad (23)$$

s.t.

$$U = E_z \left\{ f_z \left(-g(i_z) + \beta U_z^F \right) + (1 - f_z) \left(-e_z + \beta E_s \left\{ U_{z,s}^N | e_z \right\} \right) \right\}, \quad (24)$$

$$-\pi_p \chi - A i_z + \beta J(U_z^F) \geq \underline{J} - \Delta \ \forall z^1, z^2 \quad (25)$$

$$E_z \left\{ -\pi_a(e_z) \chi + \beta E_s \left\{ J(U_{z,s}^N) | e_z \right\} | z^1 \right\} \geq \underline{J} - \Delta \ \forall z^1 \quad (26)$$

$$\beta \left(E_s \left\{ U_{z,s}^N | e_z = \eta \right\} - E_s \left\{ U_{z,s}^N | e_z = 0 \right\} \right) \geq e_z \ \forall z^1, z^2 \quad (27)$$

$$J(U_z^F), J(U_{z,s}^N) \geq \underline{J} - \Delta \ \forall z^1, z^2, s \quad (28)$$

$$U_z^F, U_{z,s}^N \geq \underline{U} \ \forall z^1, z^2, s \quad (29)$$

$$f_z \in [0, 1], i_z \geq 0, \text{ and } e_z = \{0, \eta\} \ \forall z^1, z^2. \quad (30)$$

The difference between (1) – (8) and the above program is that (3), (4), and (6) have been replaced with (25), (26), and (28), respectively, for some $\Delta \geq 0$. It is clear that in this situation, our model corresponds to $\Delta = 0$ and the case of full commitment by the principal allows Δ to be

arbitrarily high. We provide below results which apply to the solution for any $\Delta \geq 0$. While all results in this section apply for any arbitrary $\Delta \geq 0$, we do not explicitly write the dependence of all of the terms on Δ .

10.3 Technical Preliminaries

We establish a set of technical results which are all proved in Appendix B which are useful for the proofs in the main text. Define the following functions

$$\mu(s) = \eta \frac{1 - \Phi(s, 0)}{\Phi(s, \eta) - \Phi(s, 0)}, \text{ and} \quad (31)$$

$$\omega(\beta, s) = \eta \frac{1/\beta - \Phi(s, 0)}{\Phi(s, \eta) - \Phi(s, 0)}. \quad (32)$$

The following lemma highlights important properties of the functions $\mu(s)$ and $\omega(\beta, s)$.

Lemma 3 *Functions $\mu(s)$ and $\omega(\beta, s)$ satisfy the following properties:*

1. $\lim_{s \rightarrow 0} \mu(s) = \infty$ and $\lim_{s \rightarrow \bar{s}} \mu(s) = \eta$,
2. $\mu'(s) < 0 \forall s \in (0, \bar{s})$,
3. $\lim_{s \rightarrow 0} \omega(\beta, s) = \lim_{s \rightarrow \bar{s}} \omega(\beta, s) = \infty$,
4. *There exists an increasing function $\hat{s}(\beta) \in [0, \bar{s}]$ which satisfies $\omega_s(\beta, s) < (>) 0$ if $s < (>) \hat{s}(\beta)$ for*

$$\frac{1/\beta - \Phi(\hat{s}(\beta), \eta)}{\Phi_s(\hat{s}(\beta), \eta)} - \frac{1/\beta - \Phi(\hat{s}(\beta), 0)}{\Phi_s(\hat{s}(\beta), 0)} = 0, \quad (33)$$

5. $\lim_{\beta \rightarrow 1} \hat{s}(\beta) = \bar{s}$, and
6. $\omega_{ss}(\beta, s) > 0$ if $s > \hat{s}(\beta)$.

The following lemma uses this characterization of $\mu(s)$ and $\omega(\beta, s)$ to establish an implication of Assumption 3. As a reminder, the exact threshold $\hat{\beta} \in (0, 1)$ in Assumption 3 is defined in Appendix B.

Lemma 4 *If $\beta > \hat{\beta}$, then*

1. *There exists $s' \in (\hat{s}(\beta), \bar{s})$ which satisfies*

$$K = \omega(\beta, s') \quad \forall K \geq g(0), \text{ and} \quad (34)$$

2. There exists $s'' \in (\widehat{s}(\beta), \bar{s})$ which satisfies

$$\left[\frac{g(i^*) - \eta}{(\pi_p - \pi_a(\eta))\chi + Ai^*} \right] (\pi_p - \pi_a(\eta))\chi + \eta = \omega(\beta, s'') \quad \text{and} \quad (35)$$

$$\beta(g(i^*) - \mu(s'')) > (g(i^*) - \eta) \left(\frac{Ai^*}{(\pi_p - \pi_a(\eta))\chi + Ai^*} \right) \quad (36)$$

for i^* defined in (9).

Finally, we can write an important implication of this lemma which is useful for the characterization of the equilibrium. Let us define $\{\underline{U}(i^*), \bar{U}(i^*), \tilde{s}(i^*), \tilde{d}(i^*)\}$ as the solution to the following system of equations for i^* defined in (9):

$$\underline{U}(i^*)(1 - \beta) = \max \left\{ -\frac{(g(i^*) - \eta)}{(\pi_p - \pi_a(\eta))\chi + Ai^*} ((\pi_p - \pi_a(\eta))\chi + \Delta(1 - \beta)) - \eta, -g(i^*) \right\}, \quad (37)$$

$$\bar{U}(i^*)(1 - \beta) = -\eta \left(\frac{1 - \Phi(\tilde{s}(i^*), 0)}{\Phi(\tilde{s}(i^*), \eta) - \Phi(\tilde{s}(i^*), 0)} \right), \quad (38)$$

$$\underline{U}(i^*)(1 - \beta) = -\eta \left(\frac{1/\beta - \Phi(\tilde{s}(i^*), 0)}{\Phi(\tilde{s}(i^*), \eta) - \Phi(\tilde{s}(i^*), 0)} \right), \quad (39)$$

$$\tilde{s}(i^*) > \widehat{s}(\beta), \text{ and} \quad (40)$$

$$\underline{U}(i^*) = -g(i^*) + \beta \left((1 - \tilde{d}(i^*)) \bar{U}(i^*) + \tilde{d}(i^*) \underline{U}(i^*) \right). \quad (41)$$

Lemma 5 *The solution to (37)–(41) exists, is unique, and admits $\tilde{d}(i^*) \in (0, 1]$ with $\tilde{d}(i^*) < 1$ if $\Delta = 0$.*

10.4 Characterization of Generalized Problem

In this section, we provide some useful lemmas for the characterization of the generalized problem. We provide some economic intuition for the lemmas, and the formal proofs are relegated to Appendix B.

Let Γ represent the set of sequential equilibrium continuation values. As a reminder, $\underline{U}$ corresponds to the lowest continuation value to the agent in this set. Furthermore, let

$$\rho^*(U) = \left\{ f_z^*(U), i_z^*(U), e_z^*(U), U_z^{F*}(U), \{U_{z,s}^{N*}(U)\}_{s \in [0, \bar{s}]} \right\}_{z \in Z}$$

correspond to a solution to the generalized problem, where it is clear that such a solution need not be unique.

Lemma 6 *Γ and $J(U)$ satisfy the following properties:*

1. Γ is convex and compact so that $J(U)$ is weakly concave, and

$$2. J(\underline{U}) = \max \{ (-\pi_p \chi - A\bar{i}) / (1 - \beta), \underline{J} - \Delta \}.$$

Lemma 6 states that $J(U)$ is concave. In addition, it characterizes $J(\underline{U})$, the welfare to the principal associated with providing the agent with the lowest credible continuation value. For the case of full commitment by the principal (i.e., Δ is arbitrarily large), it is the case that the lowest continuation to the agent is associated with intervention with maximal intensity forever so that the principal receives $(-\pi_p \chi - A\bar{i}) / (1 - \beta)$. In the case of no commitment by the principal (i.e., $\Delta = 0$), the lowest continuation value to the agent is associated with the principal receiving his min-max, which equals $\underline{J}$, since any punishment for the agent below this value could not be credibly inflicted by the principal.

Lemma 7 *Given i^* defined in (9), $J(U)$ satisfies*

$$J(U) \leq J^{\max}(U) \quad (42)$$

where

$$J^{\max}(U) = \begin{cases} \frac{-\pi_p \chi - Ag^{-1}(-U(1 - \beta))}{1 - \beta} & \text{if } U \leq \frac{-g(i^*)}{1 - \beta} \\ \frac{\psi(U)(-\pi_p \chi - Ai^*) + (1 - \psi(U))(-\pi_a(\eta)\chi)}{1 - \beta} & \text{if } U \in \left[\frac{-g(i^*)}{1 - \beta}, \frac{-\eta}{1 - \beta} \right] \end{cases} \quad (43)$$

for

$$\psi(U) = \frac{-\eta - U(1 - \beta)}{-\eta + g(i^*)}. \quad (44)$$

Lemma 8 *If (42) binds for some $U \leq -\eta / (1 - \beta)$, then*

1. *If $U \geq -g(i^*) / (1 - \beta)$,*

$$\begin{aligned} e_z^*(U) &= \eta \text{ and } U_{zs}^{N*}(U) \in [-g(i^*) / (1 - \beta), -\eta / (1 - \beta)] \quad \forall s \text{ if } f_z^*(U) = 0, \text{ and} \\ i_z^*(U) &= i^* \text{ and } U_z^{F*}(U) \in [-g(i^*) / (1 - \beta), -\eta / (1 - \beta)] \text{ if } f_z^*(U) = 1, \end{aligned}$$

2. *If $U < -g(i^*) / (1 - \beta)$, then $f_z^*(U) = 1 \quad \forall z$, $i_z^*(U) = g^{-1}(-U(1 - \beta))$ and $U_z^{F*}(U) = U$, and*

3. *$J(U_{zs}^{N*}(U)) = J^{\max}(U_{zs}^{N*}(U))$ and $J(U_z^{F*}(U)) = J^{\max}(U_z^{F*}(U)) \quad \forall s, z$.*

Lemma 7 characterizes $J^{\max}(U)$ which represents the maximal feasible welfare the principal could achieve conditional on providing the agent some continuation value below $-\eta / (1 - \beta)$, which is the welfare associated with exerting high effort forever. Lemma 8 describes the set of actions and future continuation values which sustain a continuation value U if it is the case that $J(U) = J^{\max}(U)$ so that the principal is able to achieve the maximal welfare. More specifically, if the continuation value to the agent U is between $-g(i^*) / (1 - \beta)$ and $-\eta / (1 - \beta)$, then

$J(U) = J^{\max}(U)$ is linear and the highest continuation value to the principal is provided by randomizing between intervention with intensity i^* and non-intervention with high effort by the agent in all periods. This implies that all continuation values chosen in the future are between $-g(i^*)/(1-\beta)$ and $-\eta/(1-\beta)$ and also yield the highest feasible welfare to the principal. In this regard, the value of i^* corresponds to the level of intensity which maximizes the principal's welfare conditional on randomizing between intervention and non-intervention with high effort by the agent. For continuation values below $-g(i^*)/(1-\beta)$, the principal intervenes forever with some intensity which exceeds i^* .

Lemma 9 *Given $\underline{U}(i^*)$ and $\overline{U}(i^*)$ defined in (37) – (41),*

1. (42) binds for all $U \in \left[J^{\max^{-1}} \left(\max \{ \underline{J} - \Delta, (-\pi_p \chi - A\bar{i}) / (1-\beta) \} \right), \overline{U}(i^*) \right]$,
2. (42) is a strict inequality for $U \in (\overline{U}(i^*), -\eta/(1-\beta)]$.

This lemma states that it is the case that $J(U) = J^{\max}(U)$ for the continuation values above the minimum level and below some threshold $\overline{U}(i^*)$. The economics behind this lemma is that the discount factor is sufficiently high that continuation values below $\overline{U}(i^*) < -\eta/(1-\beta)$ can be provided for the agent as efficiently as possible. This is because the principal has sufficient incentives to exert intensity weakly above i^* and the agent has sufficient incentives to exert high effort. Continuation values above $\overline{U}(i^*)$ cannot be provided as efficiently since the agent requires inducement for providing high effort which means that transitions to future periods of intervention must occur with sufficiently high probability, which lowers today's continuation value for both the principal and the agent. This keeps $\overline{U}(i^*)$ bounded away from $-\eta/(1-\beta)$, the continuation value associated with the agent exerting high effort forever.

Lemma 10 *$\overline{U}(i^*)$ satisfies the following properties:*

1. If $U \geq \overline{U}(i^*)$, then $f_z^*(U) = 0 \ \forall z$, and
2. If $U \geq \overline{U}(i^*)$ and $U_{zs}^{N*}(U) \leq \overline{U}(i^*)$ for some s, z , then

$$U_{zs}^{N*}(U) = \overline{U}(i^*) \text{ or } U_{zs}^{N*}(U) \leq \underline{U}(i^*). \quad (45)$$

The first part of Lemma 10 states that continuation values above $\overline{U}(i^*)$ are associated with zero probability of intervention today. The intuition for this is that if the agent is receiving a sufficiently high continuation value, there is no need for the principal to intervene, since intervention harms both the principal and the agent. The second part of Lemma 10 states that if $U_t \geq \overline{U}(i^*)$ and $U_{t+1} \leq \overline{U}(i^*)$, then $U_{t+1} \notin (\underline{U}(i^*), \overline{U}(i^*))$. This argument is a consequence of the MLRP property. To see a heuristic proof for this argument, suppose $J(U)$ is differentiable everywhere, and that $f_z^*(U) = 0$ and $e_z^*(U) = \eta \ \forall z$ with $U_{z,s}^{N*}(U) = U_s^{N*}(U)$ so that it is independent of z . Let λ correspond to the Lagrange multiplier on constraint (24)

and let v represent the Lagrange multiplier on constraint (27) which is assumed to bind so that $v > 0$. Suppose that it is the case that the solution admits $U_{s'}^{N*}(U) \in (\underline{U}(i^*), \overline{U}(i^*))$ and $U_{s''}^{N*}(U) \in (\underline{U}(i^*), \overline{U}(i^*))$ for some $s' \neq s''$, so that the linearity of $J(U)$ in the range $(\underline{U}(i^*), \overline{U}(i^*))$ implies that

$$J'(U_{s'}^{N*}(U)) = J'(U_{s''}^{N*}(U)) > 0. \quad (46)$$

First order conditions with respect to $U_s^{N*}(U)$ yield

$$(J'(U_{s'}^{N*}(U)) + \lambda) + v \left(1 - \frac{\Phi_s(s', 0)}{\Phi_s(s', \eta)} \right) = (J'(U_{s''}^{N*}(U)) + \lambda) + v \left(1 - \frac{\Phi_s(s'', 0)}{\Phi_s(s'', \eta)} \right), \quad (47)$$

but given (46), (47) violates the MLRP property since $\Phi_s(s, 0)/\Phi_s(s, \eta)$ is rising in s . Therefore, it cannot be that $U_s^{N*}(U) \in (\underline{U}(i^*), \overline{U}(i^*))$ with any positive probability.

10.5 Proofs of Section 4

10.5.1 Proof of Proposition 1

To establish part 1, suppose by contradiction that $U_t < 0$ and that $\Pr\{f_{t+k} = 1\} = 0 \forall k \geq 0$. Consider the payoff to the agent from a feasible strategy of choosing $e_{t+k} = 0 \forall k \geq 0$. Since $\Pr\{f_{t+k} = 1\} = 0 \forall k \geq 0$, the continuation value for the agent from following such a strategy is 0, making him strictly better off. Therefore, it is not possible for $U_t < 0$, which is a contradiction.

To establish parts 2 and 3, note that the full commitment solution corresponds to the solution to the generalized problem for which Δ is arbitrarily large so that $\underline{U}(i^*) = -g(i^*)/(1 - \beta)$ for $\underline{U}(i^*)$ defined in (37). Note that if $f_t = 1$, it must be that $U_t < \overline{U}(i^*)$ by part 1 of Lemma 10. Now consider the continuation value U_0 which the agent receives in equilibrium starting from $t = 0$. It must be that $U_0 \geq \overline{U}(i^*)$ for $\overline{U}(i^*)$ defined in (38) since if it were not the case, then Lemmas 7 and 9 imply that it is possible to increase U_0 while making both the principal and the agent strictly better off since $J(U)$ is upward sloping for $U < \overline{U}(i^*)$. By part 2 of Lemma 10,

$$\Pr\{U_{t+k} \in (\underline{U}(i^*), \overline{U}(i^*)) | U_t \geq \overline{U}(i^*)\} = 0 \forall k \geq 0,$$

which given that $U_0 \geq \overline{U}(i^*)$ implies that $\Pr\{U_t \in (\underline{U}(i^*), \overline{U}(i^*))\} = 0 \forall t$. This together with part 1 of Lemma 10 means that if $f_t = 1$, then it must be that $U_t \leq \underline{U}(i^*)$. Lemmas 8 and 9 further imply that if $f_t = 1$ and $U_t \leq \underline{U}(i^*)$ then $f_{t+k} = 1 \forall k \geq 0$, which establishes part 2. Lemmas 8 and 9 also imply that if $f_t = 1$ and $U_t \leq \underline{U}(i^*)$ then $i_{t+k} = g^{-1}(-U_t(1 - \beta)) \forall k \geq 0$, which establishes part 3. ■

10.5.2 Proof of Corollary 1

To prove this corollary it is sufficient to prove that $\lim_{t \rightarrow \infty} \Pr\{U_t \in (\underline{U}(i^*), 0)\} = 0$. This is because if $U_t = 0$, then the unique equilibrium given feasible payoffs involves $f_{t+k} = e_{t+k} = 0$

$\forall k \geq 0$. If instead $U_t \leq \underline{U}(i^*)$, then Lemmas 8 and 9 given that Δ is arbitrarily large imply that $f_{t+k} = 1$ and $i_{t+k} = g^{-1}(-U_t(1-\beta)) \forall k \geq 0$. Now suppose by contradiction that $\lim_{t \rightarrow \infty} \Pr\{U_t \in (\underline{U}(i^*), 0)\} > 0$. The proof of part 2 of Proposition 1 establishes that $\Pr\{U_t \in (\underline{U}(i^*), \overline{U}(i^*))\} = 0$ since $U_0 \geq \overline{U}(i^*)$. Therefore, it must be that $\lim_{t \rightarrow \infty} \Pr\{U_t \in [\overline{U}(i^*), 0)\} > 0$. If this is true, then it must be that

$$\lim_{t \rightarrow \infty} \Pr\{U_{t+k} \leq \underline{U}(i^*) | U_t \in [\overline{U}(i^*), 0)\} = 0 \forall k \geq 0$$

since $\Pr\{U_{t+k'} \leq \underline{U}(i^*) | U_{t+k} \leq \underline{U}(i^*)\} = 1 \forall k' \geq k \geq 0$, which would thus contradict the fact that $\lim_{t \rightarrow \infty} \Pr\{U_t \in [\overline{U}(i^*), 0)\} > 0$. Analogous arguments imply that

$$\lim_{t \rightarrow \infty} \Pr\{U_{t+k} = 0 | U_t \in [\overline{U}(i^*), 0)\} = 0 \forall k \geq 0$$

since $\Pr\{U_{t+k'} = 0 | U_{t+k} = 0\} = 1 \forall k' \geq k \geq 0$. Therefore,

$$\Pr\{U_{t+k} \in [\overline{U}(i^*), 0) | U_t \in [\overline{U}(i^*), 0)\} = 1 \forall k \geq 0.$$

However, from part 1 of Lemma 10, this means that $\Pr\{f_{t+k} = 0 | U_t \in [\overline{U}(i^*), 0)\} = 1 \forall k \geq 0$, violating part 1 of Proposition 1. Therefore, $\lim_{t \rightarrow \infty} \Pr\{U_t \in (\underline{U}(i^*), 0)\} = 0$. ■

10.6 Proofs of Section 5.1

10.6.1 Proof of Proposition 2

Conditional $U_t < 0$, part 1 follows by the same arguments as in the proof of part 1 of Proposition 1. We are left to show that it is not possible for $U_t \geq 0$. From feasibility, it is not possible that $U_t > 0$. Suppose it were the case that $U_t = 0$. Then the unique solution given feasibility would involve $f_{t+k} = e_{t+k} = 0 \forall k \geq 0$. However, if this is true, then the principal's continuation value at all dates is $-\pi_a(0)\chi/(1-\beta)$, which given Assumption 2 violates (6).

To prove parts 2 and 3, note that the case of limited commitment corresponds to a special case of the generalized problem with $\Delta = 0$. In this situation, $\underline{U}(i^*) > -g(i^*)/(1-\beta)$ for $\underline{U}(i^*)$ defined in (37). By Lemmas 7 and 9, $U_t \geq \underline{U}(i^*) \forall t$ in order to satisfy (6) since (37) implies that $J(\underline{U}(i^*)) = \underline{J}$ given $\Delta = 0$. Part 1 of Lemma 10 implies that if $f_t = 1$, then $U_t < \overline{U}(i^*)$ for $\overline{U}(i^*)$ defined in (38). Therefore, if $f_t = 1$, it must be that $U_t \in [\underline{U}(i^*), \overline{U}(i^*))$. Lemmas 8 and 9 imply that if $f_t = 1$ and $U_t \in [\underline{U}(i^*), \overline{U}(i^*))$, then $i_t = i^*$, which proves part 3. To prove part 2, suppose by contradiction if it were the case that $\Pr\{f_{t+k} = 1 | f_t = 1\} = 1 \forall k \geq 0$. This would imply given part 3 that $\Pr\{i_{t+k} = i^* > 0 | f_t = 1\} = 1 \forall k \geq 0$, so that the principal's continuation value conditional on $f_t = 1$ is $(-\pi_p\chi - Ai^*)/(1-\beta) < \underline{J}$, violating (6). ■

10.6.2 Proof of Corollary 2

Part 1 of Proposition 2 together with Lemma 10 imply that $\Pr \{U_t \leq \bar{U}(i^*) \text{ for some } t\} = 1$. The proof of parts 2 and 3 of Proposition 2 implies that

$$\Pr \{U_t \in [\underline{U}(i^*), \bar{U}(i^*)] | U_t \leq \bar{U}(i^*)\} = 1.$$

Lemmas 8 and 9 imply that if $U_t \in [\underline{U}(i^*), \bar{U}(i^*)]$, then $J(U_{t+k}) = J^{\max}(U_{t+k})$ for all $k \geq 0$. Therefore, from Lemma 8, this implies that

$$\Pr \{U_{t+k} \in [\underline{U}(i^*), \bar{U}(i^*)] | U_t \in [\underline{U}(i^*), \bar{U}(i^*)]\} = 1 \quad \forall k \geq 0.$$

This means that

$$\lim_{t \rightarrow \infty} \Pr \{U_t \in [\underline{U}(i^*), \bar{U}(i^*)]\} = 1. \quad (48)$$

From Lemmas 8 and 9, this means that in the long run, it is the case that either $f_t = 1$ and $i_t = i^*$ or $f_t = 0$ and $e_t = \eta$. ■

10.7 Proofs of Section 5.2

10.7.1 Proof of Proposition 3

By the arguments in the proof of Proposition 2, if $f_{t-k} = 1$ for some $k \geq 0$, then $U_{t-k} \in [\underline{U}(i^*), \bar{U}(i^*)]$. By the arguments of Corollary 2, this implies that

$$\Pr \{U_t \in [\underline{U}(i^*), \bar{U}(i^*)] | f_{t-k} = 1\} = 1 \quad \forall k \geq 0.$$

Given Lemmas 7 and 9, $J(\cdot)$ is linear in the range $[\underline{U}(i^*), \bar{U}(i^*)]$ with $\underline{U}(i^*)$ and $\bar{U}(i^*)$ representing the two extreme points of this linear portion. Therefore, in the equilibrium which satisfies the *Bang-Bang* property, the agent's continuation value following the realization of z_t^1 is either $\bar{U}(i^*)$ or $\underline{U}(i^*)$.

We now construct an equilibrium which satisfies this property. Suppose that the cooperative and punishment phases occur as described in the statement of the proof of Proposition 3. Let $\tilde{s}(i^*)$ and $\tilde{d}(i^*)$ correspond to the values which satisfies (37) – (41), where Lemma 5 guarantees that these values exist and satisfy $\tilde{d}(i^*) \in (0, 1]$.

We can show that this conjectured equilibrium provides welfare $J(\bar{U}(i^*))$ and $\bar{U}(i^*)$ to the principal and to the agent, respectively, during the cooperative phase, and welfare $J(\underline{U}(i^*)) = \underline{J}$ and $\underline{U}(i^*)$ to the principal and to the agent, respectively, during the punishment phase. To see why, let U^C and U^P correspond to the agent's continuation values in the cooperative and punishment phase, respectively, and define J^C and J^P for the principal analogously. Given the description of the equilibrium, these continuation values must satisfy

$$U^C = -\eta + \beta (\Phi(\tilde{s}(i^*), \eta) U^C + (1 - \Phi(\tilde{s}(i^*), \eta)) U^P), \quad (49)$$

$$U^P = -g(i^*) + \beta \left((1 - \tilde{d}(i^*)) U^C + \tilde{d}(i^*) U^P \right), \quad (50)$$

$$J^C = -\pi_a(\eta) \chi + \beta (\Phi(\tilde{s}(i^*), \eta) J^C + (1 - \Phi(\tilde{s}(i^*), \eta)) J^P), \text{ and} \quad (51)$$

$$J^P = -\pi_p \chi - Ai^* + \beta \left((1 - \tilde{d}(i^*)) J^C + \tilde{d}(i^*) J^P \right). \quad (52)$$

Given $\tilde{s}(i^*)$ and $\tilde{d}(i^*)$, by some algebra, combination of (49) and (50) implies that $U^C = \bar{U}(i^*)$ for $\bar{U}(i^*)$ satisfying (38) and $U^P = \underline{U}(i^*)$ for $\underline{U}(i^*)$ satisfying (39). Analogously, by some algebra, combination of (51) and (52) given the formula for $J(\cdot)$ implied by Lemmas 7 and 9 means that $J^C = J(\bar{U}(i^*))$ and $J^P = J(\underline{U}(i^*))$.

We now show that this conjectured equilibrium satisfies all incentive compatibility constraints. Since $J(\bar{U}(i^*)) > J(\underline{U}(i^*)) = \underline{J}$, the principal's incentive compatibility constraint is satisfied in both phases. The agent's incentive compatibility constraint need only be verified during the cooperative phase in which he exerts high effort. In this case, the values of $\bar{U}(i^*)$ and $\underline{U}(i^*)$ implied by (38) and (39) imply equation (13) so that the agent's incentive compatibility constraint is satisfied.

We have established that the *Bang-Bang* equilibrium described in the statement of Proposition 3 is an efficient sequential equilibrium. The following lemma which is proved in Appendix B shows that this equilibrium constitutes the unique efficient equilibrium satisfying the *Bang-Bang* property.

Lemma 11 *If intervention has occurred before t (i.e., $f_{t-k} = 1$ for some $k \geq 0$), then the equilibrium described in Proposition 3 is the unique efficient equilibrium which satisfies the Bang-Bang property.*

10.7.2 Proof of Proposition 4

We establish parts 2 to 4 first which allows us to establish a preliminary result which aids in the proof of part 1 which is relegated to the end.

Part 2. Equations (14) – (16) imply that given i ,

$$\underline{J}(1 - \beta) = \gamma_p(i) (-\pi_p \chi - Ai) + (1 - \gamma_p(i)) (-\pi_a(\eta) \chi) \quad (53)$$

for some $\gamma_p(i) \in [0, 1]$, since the average flow payoff to the principal under punishment must be between $-\pi_p \chi - Ai$ and $-\pi_a(\eta) \chi$, the flow payoff from punishment and cooperation, respectively. Since $\underline{J} = -\pi_p \chi / (1 - \beta)$, equation (53) implies that $\gamma_p(i)$ satisfies

$$\gamma_p(i) = \frac{(\pi_p - \pi_a(\eta)) \chi}{(\pi_p - \pi_a(\eta)) \chi + Ai}. \quad (54)$$

Moreover, equations (10) and (11) imply that

$$\underline{U}(i)(1 - \beta) = \gamma_p(i)(-g(i)) + (1 - \gamma_p(i))(-\eta), \quad (55)$$

since the average flow payoff to the agent under punishment must be between $-g(i)$ and $-\eta$, where the same weights apply to each phase as for the principal. Substituting (54) into (55), we achieve

$$\underline{U}(i)(1 - \beta) = - \left[\frac{g(i) - \eta}{(\pi_p - \pi_a(\eta))\chi + Ai} \right] (\pi_p - \pi_a(\eta))\chi - \eta. \quad (56)$$

It can be shown that the right hand side of (56) is decreasing in i for $i < i^*$ and increasing in i for $i > i^*$, so that it reaches a minimum at $i = i^*$. This is because the derivative of this function has the same sign as

$$-g'(i)((\pi_p - \pi_a(\eta))\chi + Ai) + A(g(i) - \eta), \quad (57)$$

so that this is clearly zero for $i = i^*$. Note that this derivative equals $-\infty$ at $i = 0$. Moreover, differentiating (57), we achieve

$$-g''(i)((\pi_p - \pi_a(\eta))\chi + Ai),$$

which is positive. Therefore, (57) must be negative for $i < i^*$ and positive for $i > i^*$, which completes the proof of part 2.

Part 3. We now prove part 3 by showing that $\tilde{s}(i)$ is increasing (decreasing) in i for $i < (>) i^*$. Given the proof of part 2, it is sufficient to establish that $\tilde{s}(i)$ decreases as $\underline{U}(i)$ increases. Equations (10) and (13) imply that $\overline{U}(i)$ can be rewritten as

$$\overline{U}(i)(1 - \beta) = -\eta \left(\frac{1 - \Phi(\tilde{s}(i), 0)}{\Phi(\tilde{s}(i), \eta) - \Phi(\tilde{s}(i), 0)} \right). \quad (58)$$

Using (13) to substitute in for $\overline{U}(i)$, this implies that

$$\underline{U}(i)(1 - \beta) = -\eta \left(\frac{1/\beta - \Phi(\tilde{s}(i), 0)}{\Phi(\tilde{s}(i), \eta) - \Phi(\tilde{s}(i), 0)} \right). \quad (59)$$

That there exists $\tilde{s}(i) \in [\hat{s}(\beta), \bar{s}]$ which solves (59) given (56) follows from part 1 of Lemma 4. This establishes that $\tilde{l}(i) \in [0, 1]$. Since $\tilde{s}(i) \geq \hat{s}(\beta)$, it follows from Lemma 3 that the right hand side of (59) decreases as $\tilde{s}(i)$ increases. Therefore, $\tilde{s}(i)$ decreases as $\underline{U}(i)$ increases, completing the proof of part 3.

Part 4. We now show that $\overline{U}(i)$ and $\overline{J}(i)$ are both increasing in $\tilde{s}(i)$, which combined with the proof of part 3 proves part 4. By Lemma 3, the right hand side of (58) increases in

$\tilde{s}(i)$, establishing that $\bar{U}(i)$ rises in $\tilde{s}(i)$. Equation (14) taking into account (16) implies that

$$\bar{J}(i)(1-\beta) = (1-\beta) \frac{-\pi_a(\eta)\chi + \beta(1-\Phi(\tilde{s}(i), \eta))\underline{J}}{1-\beta\Phi(\tilde{s}(i), \eta)} \quad (60)$$

The derivative of the right hand side of (60) with respect to $\tilde{s}(i)$ is equal to

$$(1-\beta) \frac{\beta\Phi_s(\tilde{s}(i), \eta)(\pi_p\chi - \pi_a(\eta)\chi)}{[1-\beta\Phi(\tilde{s}(i), \eta)]^2} > 0$$

where we have used the fact that $\Phi_s(\tilde{s}(i), \eta) > 0$ and $\pi_p > \pi_a(\eta)$. Therefore, $\bar{J}(i)$ increases if and only if $\tilde{s}(i)$ increases.

Part 1. To prove part 1, we establish the following preliminary result which is proved in Appendix B.

Lemma 12 $\bar{J}''(i) < 0$ if $i \leq i^*$.

Note that equations (15) and (16) imply that $\tilde{d}(0) = 1$. Moreover, from Lemma 5, $\tilde{d}(i^*) \in (0, 1)$ since $\Delta = 0$ in the generalized problem. Define $\bar{i} > i^*$ as follows. If $A\bar{i} \leq \beta(J(\bar{i}) - \underline{J})$, then $\bar{i} = \bar{i}$. Alternatively, if $A\bar{i} > \beta(J(\bar{i}) - \underline{J})$, then define $\bar{i}$ as the solution to $A\bar{i} = \beta(J(\bar{i}) - \underline{J})$, where this solution exists since $Ai^* < \beta(J(i^*) - \underline{J})$ given (15), (16), and $\tilde{d}(i^*) < 1$, and from part 4 which establishes that $\bar{J}'(i) < 0$ for $i > i^*$. Given (15) and (16), it is the case that $\tilde{d}(\bar{i}) \geq 0$.

We now show that $\tilde{d}(i)$ is monotonically declining in i for $i \in [0, \bar{i}]$. Note that given that $\tilde{d}(0) = 1$ and $\tilde{d}(\bar{i}) \geq 0$, this would imply that $\tilde{d}(i) \in [0, 1] \forall i \in [0, \bar{i}]$. Implicit differentiation of (15) with respect to i , given (16), yields:

$$(\bar{J}(i) - \underline{J})\tilde{d}'(i) = -A/\beta + (1 - \tilde{d}(i))\bar{J}'(i). \quad (61)$$

Note that $\tilde{d}'(i)$ has the same sign as the right hand side of (61) because $\bar{J}(i) > \underline{J}$ for $i \in [0, \bar{i}]$. To see why this is the case, note that given part 4, it is sufficient to verify that $\bar{J}(\bar{i}) > \underline{J}$ and $\bar{J}(0) > \underline{J}$ to establish this fact. That $\bar{J}(\bar{i}) > \underline{J}$ follows from the definition of $\bar{i}$. That $\bar{J}(0) > \underline{J}$ follows from the fact that $\tilde{s}(0)$ exists from the proof of part 2 and that (60) strictly exceeds $\underline{J}$. Therefore, we can focus on the right hand side of (61).

Suppose that $i > i^*$. From part 4, $\bar{J}'(i) < 0$, which means that the right hand side of (61) is negative if $\tilde{d}(i) < 1$. Since $\tilde{d}(i^*) < 1$ and $\tilde{d}(i)$ is continuous, it follows that the right hand side of (61) is negative for $i > i^*$ arbitrarily close to i^* so that $\tilde{d}(i) < 1$ for $i > i^*$ arbitrarily close to i^* . By continuity, this means that $\tilde{d}'(i)$ is strictly decreasing in the range between i^* and $\bar{i}$. Since $\tilde{d}(i^*) \in (0, 1)$ and $\tilde{d}(\bar{i}) \geq 0$. It follows that $\tilde{d}(i) \in (0, 1)$ for $i \in (i^*, \bar{i})$.

Now suppose that $i < i^*$. To prove that $\tilde{d}'(i) < 0$ in this range note that from (61), $\tilde{d}'(0) < 0$ since $\tilde{d}(0) = 1$. We now establish that the right hand side of (61) is declining in i for $i < i^*$, which

together with the fact that $\tilde{d}'(0) < 0$ means that $\tilde{d}'(i) < 0$ for $i < i^*$. Implicitly differentiating the right hand side of (61) with respect to i yields

$$\left(1 - \tilde{d}(i)\right) \bar{J}''(i) - \tilde{d}'(i) J'(i). \quad (62)$$

Suppose that $\tilde{d}'(i) \geq 0$ for some i and let $\tilde{i}$ correspond to the lowest such value of i for which this is true so that by continuity, $\tilde{d}'(\tilde{i}) = 0$. The fact that $\tilde{d}'(0) < 0$ implies that $\tilde{i} > 0$. Since $\tilde{d}'(i) < 0$ for $i < \tilde{i}$ and $\tilde{d}'(\tilde{i}) = 0$, it follows that (62) must be positive for $i = \tilde{i}$ by the continuous differentiability of $\tilde{d}'(i)$. Given that $\tilde{d}'(\tilde{i}) = 0$, this means that $\left(1 - \tilde{d}(\tilde{i})\right) \bar{J}''(\tilde{i}) > 0$. Since $\tilde{d}(0) = 1$ and $\tilde{d}'(i) < 0$ for $i < \tilde{i}$, it follows that $1 - \tilde{d}(i) > 0$. However, from Lemma 12, $J''(i) < 0$ for $i < i^*$ which means that (62) is negative for $i = \tilde{i}$, which is a contradiction. Therefore, $\tilde{d}'(i) < 0$ for $i \in (0, i^*)$. Since $\tilde{d}(0) = 1$ and $\tilde{d}(i^*) \in (0, 1)$ it follows that $\tilde{d}(i) \in (0, 1)$ for $i \in (0, i^*)$. ■

10.8 Proofs of Section 5.3

10.8.1 Proof of Proposition 5

To see how each factor affects i^* , one can implicitly differentiate (9) taking into account the concavity of $g(\cdot)$ and easily achieve these comparative statics.

To see how each factor affects $\tilde{l}(i^*)$, note that $\tilde{l}(i^*)$ is increasing in $\underline{U}(i^*)/\eta$ from the proof of part 3 of Proposition 4. Therefore, we need to show that $\underline{U}(i^*)/\eta$ for $\underline{U}(i^*)$ defined in (56) given i^* defined in (9) is increasing in A , decreasing in χ , and increasing in η . From part 3 of Proposition 4, $\underline{U}'(i^*) = 0$, which implies that it is sufficient to check the partial derivative of the right hand side of (56) with respect to A , χ , and η . It is clear by inspection that the right hand side of (56) divided by η and holding i fixed is increasing in A , decreasing in χ , and increasing in η .

To see how an increase in η affects $\tilde{d}(i^*)$ note that since this creates an increase in $\tilde{l}(i^*)$, it follows from equation (14) – (16) that $\bar{J}(i^*)$ must decline. Since i^* rises whereas $\bar{J}(i^*)$ declines, equations (15) and (16) imply that $\tilde{d}(i^*)$ declines. ■

10.9 Proofs of Section 5.4

10.9.1 Proof of Lemma 1

From Lemma 10, if $f_z^*(U) = 1$, then $U < \bar{U}(i^*)$. From Lemma 9, $J(U) = J^{\max}(U)$ if $U < \bar{U}(i^*)$ so that Lemma 8 applies. From Lemma 8, it is clear that for any Δ , if $f_z^*(U) = 1$, then $i_z^*(U) \geq i^*$. ■

10.9.2 Proof of Proposition 6

The first statement in part 1 is implied by (16) which reduces to $\underline{J}(i^*)(1 - \beta) = -\pi_p\chi$ so that it is constant. The second statement in part 1 is implied by (56) which implies that $\underline{U}(i^*)(1 - \beta)$ is constant and independent of β .

To establish part 2, consider first the value of $\overline{U}(i^*)(1 - \beta)$. (58) and (59) implicitly define $\overline{U}(i^*)$ and $\tilde{s}(i^*) \geq \hat{s}(\beta)$ as a function of β . Since $\underline{U}(i^*)(1 - \beta)$ is independent of β , the right hand side of (59) must take on the same value for all β . Since the right hand side of (59) is increasing in β and decreasing in $\tilde{s}(i^*)$ by Lemma 3, it follows that $\tilde{s}(i^*)$ is increasing in β . From Lemma 3, this means that $\overline{U}(i^*)(1 - \beta)$ defined in (58) is increasing in $\tilde{s}(i^*)$. From Lemma 3, $\lim_{\beta \rightarrow 1} \hat{s}(\beta) = \bar{s}$, which given that $\tilde{s}(i^*) \geq \hat{s}(\beta)$ implies that $\tilde{s}(i^*)$ approaches $\bar{s}$ as β approaches 1. From Lemma 3, the right hand side of (58) which equals $\overline{U}(i^*)(1 - \beta)$ approaches $-\eta$ as β approaches 1.

Now consider the value of $\overline{J}(i^*)(1 - \beta)$. Equations (10) and (11) imply that

$$\overline{U}(i^*)(1 - \beta) = \gamma_a(i^*)(-\eta) + (1 - \gamma_a(i^*))(-g(i^*)), \quad (63)$$

for some $\gamma_a(i^*) \in [0, 1]$, since the average flow payoff to the agent under cooperation must be between $-g(i^*)$ and $-\eta$, the flow payoff from punishment and cooperation, respectively. Since $\overline{U}(i^*)(1 - \beta)$ is rising in β , it follows that $\gamma_a(i^*)$ is rising in β . Moreover, equations (14) and (15) imply that

$$\overline{J}(i^*)(1 - \beta) = \gamma_a(i^*)(-\pi_a(\eta)\chi) + (1 - \gamma_a(i^*))(-\pi_p\chi - Ai^*), \quad (64)$$

which given that $\gamma_a(i^*)$ is rising in β means that $\overline{J}(i^*)(1 - \beta)$ is also rising in β . Finally since $\overline{U}(i^*)(1 - \beta)$ approaches $-\eta$ as β approaches 1, it follows from (63) that $\gamma_a(i^*)$ approaches 1 as β approaches 1 so that from (64) $\overline{J}(i^*)(1 - \beta)$ approaches $-\pi_a(\eta)\chi$ as β approaches 1. ■

11 Appendix B for “The Political Economy of Indirect Control”— Not for Publication

11.1 Definition of $\widehat{\beta}$ in Assumption 3

In this section, we provide greater detail regarding the lower bound $\widehat{\beta}$ for the discount factor in Assumption 3. $\widehat{\beta}$ is defined as

$$\widehat{\beta} = \max \left\{ \widehat{\beta}', \widehat{\beta}'' \right\} \text{ for some } \widehat{\beta}' \in (0, 1) \text{ and } \widehat{\beta}'' \in (0, 1).$$

We define the thresholds $\widehat{\beta}'$ and $\widehat{\beta}''$ below.

11.1.1 Definition of $\widehat{\beta}'$

Define $\widehat{\beta}'$ as the solution to:

$$g(0) = \omega \left(\widehat{\beta}', \widehat{s} \left(\widehat{\beta}' \right) \right) \quad (\text{A-1})$$

for $\omega(\cdot)$ defined in (32) and $\widehat{s}(\widehat{\beta}')$ defined in Lemma 3. To see that $\widehat{\beta}'$ is uniquely defined, note that the right hand side of (A-1) is strictly decreasing in $\widehat{\beta}'$. This is because by definition of $\widehat{s}(\widehat{\beta}')$,

$$\omega \left(\widehat{\beta}', \widehat{s} \left(\widehat{\beta}' \right) \right) = \min_{s \in [0, \bar{s}]} \left\{ \omega \left(\widehat{\beta}', s \right) \right\},$$

and since $\omega(\widehat{\beta}', s)$ is decreasing in $\widehat{\beta}'$, $\omega(\widehat{\beta}', \widehat{s}(\widehat{\beta}'))$ must be decreasing in $\widehat{\beta}'$ by the envelope condition. This means that if a solution to (A-1) exists, it is unique. We are left to ensure that a solution exists. The right hand side of (A-1) approaches ∞ as $\widehat{\beta}'$ approaches 0 since $\Phi(\cdot)$ is bounded between 0 and 1. Now consider the limit of the right hand side of (A-1) as $\widehat{\beta}'$ approaches 1. Substitution of (33) into (32) implies that

$$\omega \left(\widehat{\beta}', \widehat{s} \left(\widehat{\beta}' \right) \right) = \eta \frac{1}{1/\widehat{\beta}' - \Phi \left(\widehat{s} \left(\widehat{\beta}' \right), \eta \right)} = \eta \frac{1}{1 - \frac{\Phi_s \left(\widehat{s}(\beta), \eta \right)}{\Phi_s \left(\widehat{s}(\beta), 0 \right)}} \quad (\text{A-2})$$

By Lemma 3, $\widehat{s}(\widehat{\beta}')$ approaches $\bar{s}$ as $\widehat{\beta}'$ approaches 1, which means from (A-2) that $\omega(\widehat{\beta}', \widehat{s}(\widehat{\beta}'))$ approaches η as $\widehat{\beta}'$ approaches 1 since $\lim_{s \rightarrow \bar{s}} \Phi_s(s, 0)/\Phi_s(s, \eta) = \infty$. Since $g(0) \in (\eta, \infty)$ by Assumption 1, a value of $\widehat{\beta}'$ satisfying (A-1) therefore exists.

11.1.2 Definition of $\widehat{\beta}''$

In order to define $\widehat{\beta}''$, define $\widehat{\beta}'''$ as the solution to

$$\left[\frac{g(i^*) - \eta}{(\pi_p - \pi_a(\eta))\chi + Ai^*} \right] (\pi_p - \pi_a(\eta))\chi + \eta = \omega\left(\widehat{\beta}''', \widehat{s}\left(\widehat{\beta}'''\right)\right), \quad (\text{A-3})$$

for i^* defined in (9). Note that the left hand side of (A-3) would equal $g(0)$ if $i^* = 0$, but i^* defined in (9) actually maximizes the left hand side of (A-3). This implies that the left hand side of (A-3) strictly exceeds the left hand side of (A-1). Therefore, by analogous reasoning as in the case of $\widehat{\beta}'$, $\widehat{\beta}'''$ exists and is unique. There are two cases to consider in defining $\widehat{\beta}''$.

Case 1 Suppose that

$$\widehat{\beta}''' \left(g(i^*) - \mu\left(\widehat{s}\left(\widehat{\beta}'''\right)\right) \right) \geq (g(i^*) - \eta) \left(\frac{Ai^*}{(\pi_p - \pi_a(\eta))\chi + Ai^*} \right) \quad (\text{A-4})$$

for i^* defined in (9). Then, define $\widehat{\beta}'' = \widehat{\beta}'''$.

Case 2 If instead (A-4) does not hold, then define $\widehat{\beta}'' > \widehat{\beta}'''$ as the solution

$$\widehat{\beta}'' \left(g(i^*) - \mu\left(s^*\left(\widehat{\beta}''\right)\right) \right) = (g(i^*) - \eta) \left(\frac{Ai^*}{(\pi_p - \pi_a(\eta))\chi + Ai^*} \right) \quad (\text{A-5})$$

for $s^*(\beta)$ which satisfies

$$s^*(\beta) \geq \widehat{s}(\beta) \text{ and} \quad (\text{A-6})$$

$$\left[\frac{g(i^*) - \eta}{(\pi_p - \pi_a(\eta))\chi + Ai^*} \right] (\pi_p - \pi_a(\eta))\chi + \eta = \omega(\beta, s^*(\beta)), \quad (\text{A-7})$$

where $\widehat{s}(\cdot)$ is defined in Lemma 3. Since $\widehat{\beta}'' > \widehat{\beta}'''$ and $\omega_\beta(\widehat{\beta}'', s) < 0$, it follows given (A-3) that a solution $s^*(\beta)$ to (A-7) exists. Moreover, by Lemma 3, there must exist two solution to (A-7), with the highest solution corresponding to some $s^*(\beta) > \widehat{s}(\widehat{\beta}'')$ for which $\omega_s(\widehat{\beta}'', s) > 0$. Since $\omega_\beta(\widehat{\beta}'', s) < 0$ and $\omega_s(\widehat{\beta}'', s) > 0$, it follows given (A-7) that $s^{*'}(\widehat{\beta}'') > 0$, which implies that the left hand side of (A-5) is strictly increasing in $\widehat{\beta}''$, where we have appealed to the fact that $\mu'(s) < 0$. This means that if a value of $\widehat{\beta}''$ which satisfies (A-5) exists, then it is unique. Since the left hand side of (A-5) is below the right hand side for $\widehat{\beta}'' = \widehat{\beta}'''$, and since the left hand side of (A-5) is strictly increasing in $\widehat{\beta}''$, we are left to guarantee that the left hand side of (A-5) exceeds the right hand side as $\widehat{\beta}''$ approaches 1. By Lemma 3, $\widehat{s}(\widehat{\beta}'') \rightarrow \bar{s}$ as $\widehat{\beta}'' \rightarrow 1$, which given that the solution to (A-7) admits some $s^*(\beta) > \widehat{s}(\widehat{\beta}'')$, it follows that $s^*(\beta) \rightarrow \bar{s}$. Since $\lim_{s \rightarrow \bar{s}} \mu(s) = \eta$ from Lemma 3, the left hand side of (A-5) approaches $g(i^*) - \eta$ —which exceeds the right hand side—as $\widehat{\beta}''$ approaches 1. Therefore, a value of $\widehat{\beta}''$ which satisfies (A-5) exists.

11.2 Proofs of Lemmas 3-12

11.2.1 Proof of Lemma 3

To establish parts 1 and 2, differentiation of $\mu(s)$ implies that $\mu'(s)$ has the same sign as

$$\frac{1 - \Phi(s, \eta)}{\Phi_s(s, \eta)} - \frac{1 - \Phi(s, 0)}{\Phi_s(s, 0)} = \int_s^{\bar{s}} \left(\frac{\Phi_s(s', \eta)}{\Phi_s(s, \eta)} - \frac{\Phi_s(s', 0)}{\Phi_s(s, 0)} \right) ds' < 0, \quad (\text{A-8})$$

where we have used the MLRP property to establish the last inequality. That $\lim_{s \rightarrow 0} \mu(s) = \infty$ follows from the fact that the numerator in (31) approaches 1 whereas the denominator approaches 0 as s approaches 0. That $\lim_{s \rightarrow \bar{s}} \mu(s) = \eta$, follows from L'Hopital's rule:

$$\lim_{s \rightarrow \bar{s}} \eta \frac{1 - \Phi(s, 0)}{\Phi(s, \eta) - \Phi(s, 0)} = \lim_{s \rightarrow \bar{s}} \eta \left(\frac{1}{1 - \Phi_s(s, \eta) / \Phi_s(s, 0)} \right) = \eta,$$

where we have used the fact that $\lim_{s \rightarrow \bar{s}} \Phi_s(s, 0) / \Phi_s(s, \eta) = \infty$.

Part 3 follows from the fact that the numerator in (32) approaches a finite positive number whereas the denominator in (32) approaches 0 as either $s \rightarrow 0$ or $s \rightarrow \bar{s}$.

To establish part 4, let us first establish the existence of $\hat{s}(\beta)$. Given part 3 and the fact that $\omega(\beta, s)$ and $\omega_s(\beta, s)$ are well defined for $s \in (0, \bar{s})$, it follows that $\omega_s(\beta, s) < 0$ for some s which is sufficiently close to 0 and $\omega_s(\beta, s) > 0$ for some s which is sufficiently close to $\bar{s}$. Differentiation of $\omega(\beta, s)$ implies that $\omega_s(\beta, s)$ has the same sign as

$$\frac{1/\beta - \Phi(s, \eta)}{\Phi_s(s, \eta)} - \frac{1/\beta - \Phi(s, 0)}{\Phi_s(s, 0)}. \quad (\text{A-9})$$

Suppose that $\omega_s(\beta, s) > 0$ for some $s = \hat{s}$. Then it must be positive $\forall s > \hat{s}$. To see why, differentiate (A-9) to achieve

$$\frac{\Phi_{ss}(s, 0)}{\Phi_s(s, 0)} \frac{1/\beta - \Phi(s, 0)}{\Phi_s(s, 0)} - \frac{\Phi_{ss}(s, \eta)}{\Phi_s(s, \eta)} \frac{1/\beta - \Phi(s, \eta)}{\Phi_s(s, \eta)}. \quad (\text{A-10})$$

Suppose that (A-9) is weakly positive. Then (A-10) is strictly positive, where this follows from the fact that $\Phi_{ss}(\cdot) < 0$ and from the MLRP property which implies that

$$\frac{\Phi_{ss}(s, 0)}{\Phi_s(s, 0)} > \frac{\Phi_{ss}(s, \eta)}{\Phi_s(s, \eta)}. \quad (\text{A-11})$$

Therefore, if $\omega_s(\beta, s) > 0$ for some $\hat{s}$, then it must be positive $\forall s > \hat{s}$. Together with part 3, this means that there exists some $\hat{s}(\beta)$ such that $\omega_s(\beta, s) < (>) 0$ if $s < (>) \hat{s}(\beta)$.

Let us now show that $\hat{s}(\beta)$ is increasing. Given its definition, $\hat{s}(\beta)$ must satisfy the following equation

$$\frac{1/\beta - \Phi(\hat{s}(\beta), \eta)}{\Phi_s(\hat{s}(\beta), \eta)} - \frac{1/\beta - \Phi(\hat{s}(\beta), 0)}{\Phi_s(\hat{s}(\beta), 0)} = 0. \quad (\text{A-12})$$

Note that

$$-\frac{\Phi(\widehat{s}(\beta), \eta)}{\Phi_s(\widehat{s}(\beta), \eta)} + \frac{\Phi(\widehat{s}(\beta), 0)}{\Phi_s(\widehat{s}(\beta), 0)} = \int_0^{\widehat{s}(\beta)} \left[-\frac{\Phi_s(s, \eta)}{\Phi_s(\widehat{s}(\beta), \eta)} + \frac{\Phi_s(s, 0)}{\Phi_s(\widehat{s}(\beta), 0)} \right] ds < 0, \quad (\text{A-13})$$

where we have used the MLRP property to establish the last inequality. Given (A-12) and (A-13), it follows that $(1/\Phi_s(\widehat{s}(\beta), \eta) - 1/\Phi_s(\widehat{s}(\beta), 0))/\beta > 0$, so that $\Phi_s(\widehat{s}(\beta), 0) > \Phi_s(\widehat{s}(\beta), \eta)$ which means that the left hand side of (A-12) is decreasing in β holding $\widehat{s}(\beta)$ fixed. Since the left hand side of (A-12) is increasing in $\widehat{s}(\beta)$ holding β fixed, it follows that $\widehat{s}(\beta)$ is an increasing function. This establishes part 4.

To establish part 5, note that by part 4 $\widehat{s}(\beta)$ is increasing in β . Suppose by contradiction that $\lim_{\beta \rightarrow 1} \widehat{s}(\beta) = \bar{s}' < \bar{s}$. Taking the limit of the left hand side of (A-12) as β approaches 1 then implies that

$$\frac{1 - \Phi(\bar{s}', \eta)}{\Phi_s(\bar{s}', \eta)} - \frac{1 - \Phi(\bar{s}', 0)}{\Phi_s(\bar{s}', 0)} = 0 \quad (\text{A-14})$$

for some $\bar{s}' < \bar{s}$. However, (A-14) contradicts (A-8). Therefore, $\lim_{\beta \rightarrow 1} \widehat{s}(\beta) = \bar{s}$. This establishes part 5.

To establish part 6, note that

$$\begin{aligned} \frac{\omega_{ss}(\beta, s)}{\eta} &= 2 \frac{\omega_s(\beta, s)}{\eta} \frac{\Phi_s(s, 0) - \Phi_s(s, \eta)}{\Phi(s, \eta) - \Phi(s, 0)} \\ &\quad + \frac{\Phi_{ss}(s, 0)(1/\beta - \Phi(s, \eta)) - \Phi_{ss}(s, \eta)(1/\beta - \Phi(s, 0))}{(\Phi(s, \eta) - \Phi(s, 0))^2} \end{aligned} \quad (\text{A-15})$$

By our previous arguments, $\Phi_s(\widehat{s}(\beta), 0) > \Phi_s(\widehat{s}(\beta), \eta)$ so that if $s \geq \widehat{s}(\beta)$, the first term on the right hand side of (A-15) is positive since $\omega_s(\beta, s) > 0$ and $\Phi_s(s, 0) - \Phi_s(s, \eta) > 0$ by the MLRP property. (A-11) and the fact that $\Phi_s(s, 0) - \Phi_s(s, \eta) > 0$ implies that $\Phi_{ss}(s, 0) > \Phi_{ss}(s, \eta)$, which together with the fact that $\Phi_{ss}(\cdot) < 0$ and $1/\beta - \Phi(s, 0) > 1/\beta - \Phi(s, \eta) > 0$ implies that the second term on the right hand side of (A-15) is also positive. Therefore, $\omega_{ss}(\beta, s) > 0$. ■

11.2.2 Proof of Lemma 4

To establish part 1, note there exists a solution s' to (34) for $K = g(0)$ if $\beta = \widehat{\beta}'$ given the definition of $\widehat{\beta}'$ in Section 11.1.1. Note that from (32), $\omega_\beta(\beta, s) < 0$. From Lemma 3, $\omega_s(\beta, s) > 0$ if $s \geq \widehat{s}(\beta)$ with $\lim_{s \rightarrow \bar{s}} \omega(\beta, s) = \infty$. This means that if $\beta > \widehat{\beta} \geq \widehat{\beta}'$, there exists a solution $s' \geq \widehat{s}(\beta)$ to (34) holding K fixed. Moreover, since $\lim_{s \rightarrow \bar{s}} \omega(\beta, s) = \infty$, if a solution exists for K' , then it exists for $K'' > K'$.

To establish part 2, note that there exists a solution s'' to (35) which satisfies (36) if $\beta = \widehat{\beta}''$ given the definition of $\widehat{\beta}''$ in Section 11.1.2. By analogous reasoning as in part 1, if $\beta > \widehat{\beta} \geq \widehat{\beta}''$, there exists a solution $s'' \geq \widehat{s}(\beta)$ to (35). We are left to verify that such a solution also satisfies (36). Suppose that $\beta > \widehat{\beta}''$. The value of s'' which satisfies (35) is increasing in β , which implies that the value of $\mu(s'')$ is declining in β since $\mu(\cdot)$ is a decreasing function by Lemma 3.

Therefore, the left hand side of (36) rises as β rises, whereas the right hand side stays constant. Since the left hand side weakly exceeds the right hand side of (36) for $\beta = \widehat{\beta}''$ by definition, (36) is satisfied for $\beta > \widehat{\beta}''$. ■

11.2.3 Proof of Lemma 5

$\underline{U}(i^*)$ is uniquely determined according to equation (37). Conditional on this value of $\underline{U}(i^*)$, equations (39) and (40) uniquely determine $\widetilde{s}(i^*)$. To see why, note that a value of $\widetilde{s}(i^*)$ which solves (39) exists by part 1 of Lemma 4 since the right hand side of (37) is strictly below $-g(0)$ given the definition of i^* in (9). That this solution is unique follows from Lemma 3 which guarantees that the right hand side of (39) is monotonic in $\widetilde{s}(i^*)$ for $\widetilde{s}(i^*) \geq \widehat{s}(\beta)$. Since $\widetilde{s}(i^*)$ is uniquely determined, this means that $\overline{U}(i^*)$ is uniquely determined by equation (38). Finally, it is clear that $\widetilde{d}(i^*)$ which solves (41) exists and is unique since $\underline{U}(i^*)$ and $\overline{U}(i^*)$ are uniquely determined.

We are left to show that $\widetilde{d}(i^*) \in (0, 1]$ with $\widetilde{d}(i^*) < 1$ if $\Delta = 0$. Given (41), we can prove this by showing that

$$\underline{U}(i^*) \leq \frac{\underline{U}(i^*) + g(i^*)}{\beta} \quad (\text{A-16})$$

which is strict if $\Delta = 0$ and

$$\frac{\underline{U}(i^*) + g(i^*)}{\beta} < \overline{U}(i^*). \quad (\text{A-17})$$

(A-16) is implied by (37) which implies that $\underline{U}(i^*)(1 - \beta) \geq -g(i^*)$ which is strict if $\Delta = 0$. To see why (A-17) holds, note that this is trivially satisfied if $\underline{U}(i^*)$ defined in (37) equals $-g(i^*) / (1 - \beta)$, which must be below $\overline{U}(i^*)$ given (38) and (39). Suppose instead that $\underline{U}(i^*)$ defined in (37) exceeds $-g(i^*) / (1 - \beta)$. To see why (A-17) holds in this case, it is sufficient to check this for $\Delta = 0$. To see why, from (37), as Δ decreases, $\underline{U}(i^*)$ increases. Moreover, from (39) and (40), as Δ decreases, then $\widetilde{s}(i^*)$ decreases, where this follows from Lemma 3. Moreover, from Lemma 3, this implies that $\overline{U}(i^*)$ defined in (38) decreases. Therefore, it is sufficient to check (A-17) $\Delta = 0$. To do this, note that the value of $\widetilde{s}(i^*)$ in this case which is determined from the combination of equations (37) and (39) also coincides with s'' in (35) in the second part of Lemma 4. By some algebra, the inequality in (36) in the second part of Lemma 4 implies that (A-17) holds. ■

11.2.4 Proof of Lemma 6

Part 1. Let Λ correspond to the set of allocations α which are feasible and satisfy (20) – (22) for $\underline{J}$ replaced by $\underline{J} - \Delta$ so that the solution corresponds to that of the generalized problem. Consider two continuation value pairs $\{V'_p, V'_a\} \in \Gamma$ and $\{V''_p, V''_a\} \in \Gamma$ with corresponding allocations α' and α'' , where an allocation is defined in (19). $\alpha'|_{z_0^1}$ corresponds to the continuation allocation

conditional on z_0^1 , and $\alpha''|_{z_0^1}$. Is defined analogously. It must be that

$$\{V_p^\kappa, V_a^\kappa\} = \{\kappa V_p' + (1 - \kappa) V_p'', \kappa V_a' + (1 - \kappa) V_a''\} \in \Gamma \quad \forall \kappa \in (0, 1).$$

Define $\alpha^\kappa = \left\{ \alpha^\kappa|_{z_0^1} \right\}_{z_0^1 \in [0, 1]}$ as follows:

$$\alpha^\kappa|_{z_0^1} = \begin{cases} \alpha'|_{z_0^1/\kappa} & \text{if } z_0^1 \in [0, \kappa) \\ \alpha''|_{(z_0^1 - \kappa)/(1 - \kappa)} & \text{if } z_0^1 \in [\kappa, 1] \end{cases},$$

where $\alpha'|_{z_0^1/\kappa}$ for $z_0^1 \in [0, \kappa)$ is identical to $\alpha'|_{z_0^1}$ with the exception that z_0^1/κ replaces z_0^1 in all information sets q_t^0 and q_t^1 , and $\alpha''|_{(z_0^1 - \kappa)/(1 - \kappa)}$ for $z_0^1 \in [\kappa, 1]$ is analogously defined. α^κ achieves $\{V_p^\kappa, V_a^\kappa\}$, and since $\alpha', \alpha'' \in \Lambda$, then $\alpha^\kappa \in \Lambda$.

Γ is bounded since $V_j(\alpha)$ is bounded for $j = p, a$. To show that Γ is closed, consider a sequence $\{V_{pn}', V_{an}'\} \in \Gamma$ such that $\lim_{n \rightarrow \infty} \{V_{pn}', V_{an}'\} = \{V_p', V_a'\}$. There exists one corresponding sequence of allocations α'_n which converges to α'_∞ since $V_j(\alpha'_n)$ is continuous in α'_n . Since every element of α'_n at q_t^0 is contained in a closed and bounded set, and since (20) – (22) are weak inequalities, then Λ is closed and $\alpha'_\infty \in \Lambda$. Since $\beta \in (0, 1)$, then by the Dominated Convergence Theorem, $V_j(\alpha'_\infty) = V_j'$ for $j = p, a$. Therefore $\{V_a', V_p'\} \in \Gamma$ so that Γ is compact.

Since Γ is a convex set, and since $J(U)$ corresponds to the highest value of V_p conditional on $V_a = U$, it follows that $J(U)$ is concave.

Part 2. To show that $J(\underline{U}) = \max \{(-\pi_p \chi - A\bar{i}) / (1 - \beta), \underline{J} - \Delta\}$, note that it is not possible that $J(\underline{U}) < \underline{J} - \Delta$ since this violates (28). Suppose that $\underline{J} - \Delta \leq (-\pi_p \chi - A\bar{i}) / (1 - \beta)$. From feasibility, $\underline{U} \geq -g(\bar{i}) / (1 - \beta)$. In this situation, $\underline{U} = -g(\bar{i}) / (1 - \beta)$, since $f_z^*(\underline{U}) = 1$, $i^*(\underline{U}) = \bar{i}$, and $U_z^{F*}(\underline{U}) = \underline{U} \quad \forall z$ satisfies (25) and (28) and provides a continuation value of $-g(\bar{i}) / (1 - \beta)$ to the agent. This solution is unique, since from feasibility any other solution provides a continuation value to the agent strictly larger than $-g(\bar{i}) / (1 - \beta)$. Therefore, $J(\underline{U}) = (-\pi_p \chi - A\bar{i}) / (1 - \beta)$. Suppose instead that $\underline{J} - \Delta > (-\pi_p \chi - A\bar{i}) / (1 - \beta)$. Suppose by contradiction that $J(\underline{U}) > \underline{J} - \Delta$. One can show that if this were true then $\underline{U}$ would not correspond to the lowest continuation value to the agent. To show this, note that by Assumptions 1 and 2 and equations (25) and (28), it must be that $f_z^*(\underline{U}) = 1 \quad \forall z$. If it were the case instead that $f_z^*(\underline{U}) = 0$ for some z , then a perturbation to $\hat{f}_z^*(\underline{U}) = 1$ and $\hat{i}_z^*(\underline{U}) = 0$ for all such z would strictly reduce $\underline{U}$ while continuing to satisfy (25) – (30). Since $f_z^*(\underline{U}) = 1 \quad \forall z$, note that if it were the case that $J(\underline{U}) > \max \{(-\pi_p \chi - A\bar{i}) / (1 - \beta), \underline{J} - \Delta\}$, then would be possible to increase $i_z^*(\underline{U})$ so as to reduce the agent's welfare while continuing to satisfy (25) – (30), contradicting the fact that $\underline{U}$ is the agent's lowest continuation value. Therefore, $J(\underline{U}) = \max \{(-\pi_p \chi - A\bar{i}) / (1 - \beta), \underline{J} - \Delta\}$. ■

11.2.5 Proof of Lemma 7

Define $J^{\max}(U)$ as follows, where α , $\mathcal{F}$, q_t^0 , and q_t^1 are defined in Section 10.1:

$$J^{\max}(U) = \max_{\alpha \in \mathcal{F}} E_0 \sum_{t=0}^{\infty} \beta^t u_p(f_t(q_t^0), i_t(q_t^1), e_t(q_t^1), s_t) \text{ s.t.} \quad (\text{A-18})$$

$$U \geq E_0 \sum_{t=0}^{\infty} \beta^t u_a(f_t(q_t^0), i_t(q_t^1), e_t(q_t^1), s_t) \quad (\text{A-19})$$

Since this program maximizes the same object as the original program but under strictly fewer constraints, it is clear that (42) must hold for some $J^{\max}(U)$ which solves (A-18) – (A-19).

We now characterize $J^{\max}(U)$ for $U \leq -\eta/(1-\beta)$ to complete the proof. We proceed in four steps.

Step 1. (A-19) must bind. To see why, suppose this were not the case so that the solution lets (A-19) remain slack. Then it cannot be that $f_t(q_t^0) = 0$ and $e_t(q_t^1) = \eta$ for all q_t^0 and q_t^1 , since if this were the case, then the right hand side of (A-19) would equal $-\eta/(1-\beta)$. However, since $U \leq -\eta/(1-\beta)$, this would mean that satisfaction of (A-19) requires (A-19) to bind, which is a contradiction. Therefore, if (A-19) is slack, then there exists some q_t^0 and q_t^1 for which it is not the case that $f_t(q_t^0) = 0$ and $e_t(q_t^1) = \eta$. Consider an alternative solution to the problem which is identical to the original solution with the exception that $f_t(q_t^0) = 0$ and $e_t(q_t^1) = \eta$ for some such q_t^0 and q_t^1 . Since $f_t(q_t^0) = 0$ and $e_t(q_t^1) = \eta$ maximizes the principal's static payoff, this alternative solution yields a strict increase in (A-18) while continuing to satisfy (A-19). Therefore, (A-19) cannot be slack.

Step 2. The solution to (A-18) – (A-19) admits $i_t(q_t^1) = i$ for some $i \geq 0$ for all q_t^1 for which $f_t(q_t^0) = 1$. To see why, suppose this were not the case. Then consider a perturbation which is identical to the original solution with the exception that if $f_t(q_t^0) = 1$, then $i_t(q_t^1)$ is replaced with

$$\hat{i}_t(q_t^1) = E_0 \sum_{k=0}^{\infty} \frac{\beta^k f_k(q_k^0) i_k(q_k^1)}{E_0 \sum_{k=0}^{\infty} \beta^k f_k(q_k^0)}.$$

This perturbation does not change the right hand side of (A-18) but it relaxes (A-19) since $g(\cdot)$ is a concave function. However, this contradicts the fact that (A-19) must bind in the optimum.

Step 3. The solution to (A-18) – (A-19) admits $e_t(q_t^1) = \eta$ for all q_t^1 for which $f_t(q_t^0) = 0$. To see why, suppose this were not the case. Then consider a perturbation which is identical to the original solution with the exception that if $f_t(q_t^0) = 0$ then $e_t(q_t^1) = \eta$. This strictly increases the right hand side of (A-18) while continuing to satisfy (A-19).

Step 4. Given steps 1-3, (A-18) – (A-19) can be rewritten as

$$J^{\max}(U)(1-\beta) = \max_{f \in [0,1], i \geq 0} \{ f(-\pi_p \chi - A i) + (1-f)(-\pi_a(\eta) \chi) \} \text{ s.t.} \quad (\text{A-20})$$

$$U(1-\beta) = f(-g(i)) + (1-f)(-\eta). \quad (\text{A-21})$$

for $f = (E_0 \sum_{t=0}^{\infty} \beta^t f_t(q_t^0)) (1 - \beta)$. If $U = -\eta / (1 - \beta)$, then the solution admits $f = 0$ since $g(0) > \eta$ so that (43) holds in this case. Now suppose that $U < -\eta / (1 - \beta)$ so that this is not possible and $f > 0$. Suppose first that the constraint that $f \leq 1$ does not bind. Then first order conditions to this program imply that the optimal value of i is i^* defined in (9). In order that $f \leq 1$ be satisfied, it must be that $U \geq -g(i^*) / (1 - \beta)$. This means that in this case the value of f which satisfies (A-21) given $i = i^*$ satisfies $f = \psi(U)$ and (43) holds in this case. If instead $U < -g(i^*) / (1 - \beta)$, then $f = 1$ and satisfaction of (A-21) implies that $i = g^{-1}(-U(1 - \beta))$, so that (43) holds in this case. ■

11.2.6 Proof of Lemma 8

Part 1. If (42) binds then $J(U)$ satisfies (A-18) – (A-19). Suppose first that $U(1 - \beta) \in [-g(i^*), -\eta]$. Then from step 4 in the proof of Lemma 7, it is necessary that for all q_t^0, q_t^1 , if $f_t(q_t^0) = 0$ then $e_t(q_t^1) = \eta$ and if $f_t(q_t^0) = 1$ then $i_t(q_t^1) = i^*$, where the value of $f = (E_0 \sum_{t=0}^{\infty} \beta^t f_t(q_t^0)) (1 - \beta)$ is unique. This implies that $e_z^*(U) = \eta$ if $f_z^*(U) = 0$ and $i_z^*(U) = i^*$ if $f_z^*(U) = 1$. Moreover, given (24), this also means that there exists some $f'_{zs} \in [0, 1]$ and $f''_z \in [0, 1]$ such that

$$U_{zs}^{N*}(U)(1 - \beta) = f'_{zs}(-g(i^*)) + (1 - f'_{zs})(-\eta) \quad \text{and} \quad (\text{A-22})$$

$$U_z^{F*}(U)(1 - \beta) = f''_z(-g(i^*)) + (1 - f''_z)(-\eta) \quad (\text{A-23})$$

for f'_{zs} and f''_z which satisfy

$$\frac{f(-g(i^*)) + (1 - f)(-\eta)}{1 - \beta} = E_z \left\{ \frac{f_z^*(U)(-g(i^*) + \beta U_z^{F*}(U)) +}{(1 - f_z^*(U))(-\eta + \beta E_s \{U_{z,s}^{N*}(U) | e_z = \eta\})} \right\}$$

given (A-22) and (A-23) and $f_z^*(U)$. This proves part 1.

Part 2. Now suppose that $U(1 - \beta) < -g(i^*)$. Then from step 4 in the proof of Lemma 7, $f_t(q_t^0) = 1$ for all q_t^0 and $i_t(q_t^1) = g^{-1}(-U(1 - \beta))$ for all q_t^1 . From (24), this means that for all z $f_z^*(U) = 1$, $i_z^*(U) = g^{-1}(-U(1 - \beta))$, and $U_z^{F*}(U) = U$. This proves part 2.

Part 3. By (42) and part 1, $J(U_{zs}^{N*}(U)) \leq J^{\max}(U_{zs}^{N*}(U))$ and $J(U_z^{F*}(U)) \leq J^{\max}(U_z^{F*}(U))$ $\forall s, z$. Suppose it were the case that these weak inequalities did not bind. Suppose first if $U \geq -g(i^*) / (1 - \beta)$ but $J(U_{zs}^{N*}(U)) < J^{\max}(U_{zs}^{N*}(U))$ or $J(U_z^{F*}(U)) < J^{\max}(U_z^{F*}(U))$. From step 4 in the proof of Lemma 7, it is necessary that for all q_t^0, q_t^1 , if $f_t(q_t^0) = 0$ then $e_t(q_t^1) = \eta$ and if $f_t(q_t^0) = 1$ then $i_t(q_t^1) = i^*$. This means by analogous arguments as those in the proof of part 1 that there exists some $f'_{zs} \in [0, 1]$ and $f''_z \in [0, 1]$ which satisfy (A-22) and (A-23) and which also satisfy

$$J(U_{zs}^{N*}(U))(1 - \beta) = f'_{zs}(-\pi_p \chi - A i^*) + (1 - f'_{zs})(-\pi_a(\eta) \chi) \quad \text{and} \quad (\text{A-24})$$

$$J(U_z^{F*}(U))(1 - \beta) = f''_z(-\pi_p \chi - A i^*) + (1 - f''_z)(-\pi_a(\eta) \chi). \quad (\text{A-25})$$

(A-24) and (A-25) together with (A-22) and (A-23) given (43) imply that $J(U_{zs}^{N*}(U)) = J^{\max}(U_{zs}^{N*}(U))$ and $J(U_z^{F*}(U)) = J^{\max}(U_z^{F*}(U))$. Suppose instead if $U < -g(i^*)/(1-\beta)$. From part 2, if $U < -g(i^*)/(1-\beta)$ then $J(U_z^{F*}(U)) = (-\pi_p\chi - Ag^{-1}(-U(1-\beta)))/(1-\beta)$ which equals $J^{\max}(U_z^{F*}(U))$ from Lemma 7. Therefore, in this case, $J(U_z^{F*}(U)) = J^{\max}(U_z^{F*}(U))$, and $J(U_{zs}^{N*}(U))$ is irrelevant since $f_z^*(U) = 1$. This proves part 3. ■

11.2.7 Proof of Lemma 9

Part 1. To establish this, we consider several cases.

Case 1. Suppose that

$$U < J^{\max^{-1}}(\max\{\underline{J} - \Delta, (-\pi_p\chi - A\bar{i})/(1-\beta)\}), \quad (\text{A-26})$$

then from the arguments in the proof of Lemma 6, U is not incentive compatible for the principal.

Case 2. Suppose that (A-26) does not hold and that

$$J^{\max^{-1}}(\max\{\underline{J} - \Delta, (-\pi_p\chi - A\bar{i})/(1-\beta)\}) < -g(i^*)/(1-\beta). \quad (\text{A-27})$$

Consider U which satisfies

$$U \in [J^{\max^{-1}}(\max\{\underline{J} - \Delta, (-\pi_p\chi - A\bar{i})/(1-\beta)\}), -g(i^*)/(1-\beta)], \quad (\text{A-28})$$

Lemma 8 characterizes the unique solution to the recursive program if (42) binds, where it must be that $f_z^*(U) = 1$, $i_z^*(U) = g^{-1}(-U(1-\beta))$ and $U_z^{F*}(U) = U$. This solution clearly satisfies (24). To check incentive compatibility, note that since

$$J^{\max}(U) = (-\pi_p\chi - Ai_z^*(U))/(1-\beta) \geq \underline{J} - \Delta,$$

(25) and (28) are satisfied so that the solution is incentive compatible for the principal. The incentive compatibility constraints on the agent are trivially satisfied since he does not choose any actions. This establishes that if $J^{\max^{-1}}(\max\{\underline{J} - \Delta, (-\pi_p\chi - A\bar{i})/(1-\beta)\}) < -g(i^*)/(1-\beta)$, then (42) binds for all continuation values in (A-28).

Case 3. We are left to study two cases simultaneously when (A-26) does not hold. First, if (A-27) holds but $U \geq -g(i^*)/(1-\beta)$. Second, if (A-27) does not hold. These two cases can be studied together since they require use to characterize the equilibrium for values of U which weakly exceed $\underline{U}(i^*)$. Specifically, in the first, case (37) and (43) imply that $\underline{U}(i^*) = -g(i^*)/(1-\beta)$. In the second case, given the definition of $J^{\max}(\cdot)$ in (43) and given the definition of $\underline{U}(i^*)$ in (37), this implies that $J^{\max}(\underline{U}(i^*)) = \underline{J} - \Delta$. Therefore, it is not possible for $U \leq \underline{U}(i^*)$, since this would imply that $J(U) \leq J^{\max}(U) < \underline{J} - \Delta$, violating (28).

Let us conjecture the following solutions $\forall z$ for $U \in [\underline{U}(i^*), \bar{U}(i^*)]$:

$$E_z f_z^*(U) = (\bar{U}(i^*) - U) / (\bar{U}(i^*) - \underline{U}(i^*)), i_z^*(U) = i^*, e_z^*(U) = \eta, \quad (\text{A-29})$$

$$U_z^{F*}(U) = (\underline{U}(i^*) + g(i^*)) / \beta, \quad (\text{A-30})$$

$$U_{zs}^{*N}(U) = \begin{cases} \bar{U}(i^*) & \text{if } s \leq \tilde{s}(i^*) \\ \underline{U}(i^*) & \text{if } s > \tilde{s}(i^*) \end{cases} \quad (\text{A-31})$$

for $\bar{U}(i^*)$ and $\tilde{s}(i^*)$ defined in (37)–(41). In the below steps, we prove that this solution satisfies (24)–(30) and generates welfare $J^{\max}(U)$ for the principal.

Step 1. $\tilde{s}(i^*)$ exists and $E_z f_z^*(U) \in [0, 1]$ by Lemma 5 which implies that $U_z^{F*}(U) \in [\underline{U}(i^*), \bar{U}(i^*)]$.

Step 2. The solution is feasible, satisfies the promise keeping constraint, and generates a continuation value to the principal equal to $J^{\max}(U)$. To see why, note that the solution is feasible since $e_z^*(U)$ and $i_z^*(U)$ are well-defined and since $U \in [\underline{U}(i^*), \bar{U}(i^*)]$ implies that $E_z f_z^*(U) \in [0, 1]$. Given the value of $E_z f_z^*(U)$, to show that the solution satisfies the promise keeping constraint, it is sufficient to show that, conditional on $f_z^*(U) = 1$, the agent receives $\underline{U}(i^*)$ and conditional on $f_z^*(U) = 0$, the agent receives $\bar{U}(i^*)$. Note that given the solution, if $f_z^*(U) = 1$, the agent receives $-g(i^*) + \beta U_z^{F*}(U) = \underline{U}(i^*)$. Moreover, if $f_z^*(U) = 0$, the agent receives

$$-\eta + \beta (\Phi(\tilde{s}(i^*), \eta) \bar{U}(i^*) + (1 - \Phi(\tilde{s}(i^*), \eta)) \underline{U}(i^*)),$$

which after the substitution of (38) and (39) implies that the agent receives $\bar{U}(i^*)$. To show that the solution generates a continuation value to the principal equal to $J^{\max}(U)$, note that since $U_z^{F*}(U) \in [\underline{U}(i^*), \bar{U}(i^*)]$ from step 1 and $U_{zs}^{*N}(U) \in [\underline{U}(i^*), \bar{U}(i^*)]$, given (A-29) $\forall U \in [\underline{U}(i^*), \bar{U}(i^*)]$, it must be that (A-21) is satisfied for $i = i^*$ and some $f \in [0, 1]$, which means that $J(U)$ equals $J^{\max}(U)$ defined in (A-20).

Step 3. The solution satisfies all incentive compatibility constraints. To see why, suppose first that $f_z^*(U) = 1$. If $J^{\max^{-1}}(\max\{\underline{J} - \Delta, (-\pi_p \chi - A\bar{i}) / (1 - \beta)\}) \leq -g(i^*) / (1 - \beta)$, then from (43), $J^{\max}(\underline{U}(i^*)) \geq \underline{J} - \Delta$, which means given (37) that $\underline{U}(i^*) = -g(i^*) / (1 - \beta)$. Therefore, from (A-30) $U_z^{F*}(U) = \underline{U}(i^*)$. This implies that

$$-\pi_p \chi - A i^* + \beta J(U_z^{F*}(U)) = (-\pi_p \chi - A i^*) / (1 - \beta) \geq \underline{J} - \Delta \quad (\text{A-32})$$

so that (25) is satisfied. Now suppose that $J^{\max^{-1}}(\max\{\underline{J} - \Delta, (-\pi_p \chi - A\bar{i}) / (1 - \beta)\}) \geq -g(i^*) / (1 - \beta)$. Note that $J(U_z^{F*}(U)) = J^{\max}(U_z^{F*}(U))$ so that (43) can be used to calculate $J(U_z^{F*}(U))$ taking into account that $U_z^{F*}(U)$ can be calculated from (37) and (A-30). By some algebra this implies that $-\pi_p \chi - A i^* + \beta J(U_z^{F*}(U)) = \underline{J} - \Delta$ so that (25) is satisfied. In both cases, (28) is satisfied since $J(U_z^{F*}(U)) \geq \underline{J} - \Delta$.

Now suppose that $f_z^*(U) = 0$. (38) and (39) imply that

$$\begin{aligned} & \beta \left(E_s \{ U_{z,s}^N | e_z = \eta \} - E_s \{ U_{z,s}^N | e_z = 0 \} \right) = \\ & \beta \left(\Phi(\tilde{s}(i^*), \eta) - \Phi(\tilde{s}(i^*), 0) \right) (\bar{U}(i^*) - \underline{U}(i^*)) = \eta \end{aligned}$$

so that the agent is indifferent between choosing $e_z = 0$ and $e_z = \eta$ so that (27) is satisfied. From (38) and (39), $\bar{U}(i^*) \geq \underline{U}(i^*) \geq \underline{U}$ so that (29) is satisfied. To check (28), note that $J(U)$ is rising in U so that it is sufficient to check that $J(\underline{U}(i^*)) \geq \underline{J} - \Delta$, and this follows from the definition of $\underline{U}(i^*)$ in (37) together with (43). Finally, (26) is satisfied since $e_z^*(U) = \eta$ and $J(U_{zs}^{N*}(U)) \geq \underline{J} - \Delta$.

Part 2. Part 1 establishes that (42) binds for

$$U \in \left[J^{\max^{-1}} \left(\max \{ \underline{J} - \Delta, (-\pi_p \chi - A\bar{i}) / (1 - \beta) \} \right), \bar{U}(i^*) \right].$$

Suppose by contradiction that there exists some $U' \in (\bar{U}(i^*), -\eta / (1 - \beta)]$ for which (42) is an equality, and let U' denote the highest such value for which this is true. Note that given weak the concavity of $J(\cdot)$ and the linearity of $J^{\max}(U)$ for U between $\underline{U}(i^*)$ and U' , this implies that if $U \in [\underline{U}(i^*), U']$, then (42) binds. We establish that it is not possible for such a $U' > \bar{U}(i^*)$ to exist in three steps.

Step 1. It is necessary that $f_z^*(U') = 0 \forall z$. To see why, suppose by contradiction that $E_z(f_z^*(U')) > 0$. Since U' is the highest value of U for which (42) binds, application of Lemma 8 implies that $e_z^*(U') = \eta$, $i_z^*(U') = i^*$, $U_z^{F*}(U') \leq U'$, and $U_{zs}^{N*}(U') \leq U' \forall z, s$. Moreover, since $U' \geq \bar{U}(i^*) > -g(i^*) / (1 - \beta)$, it follows that $-g(i^*) + \beta U_z^{F*}(U') < U'$, so that if $E_z(f_z^*(U')) > 0$, then

$$E_z \{ -g(i^*) + \beta U_z^{F*}(U') \} < U'.$$

Thus, for (24) to hold it is necessary that

$$E_z \{ -\eta + \beta E_s \{ U_{z,s}^{N*}(U') | e_z = \eta \} \} > U'. \quad (\text{A-33})$$

However, if this is the case, then there exists a value of $U'' > U'$ equal to the left hand side of (A-33) for which $J(U'') = J^{\max}(U'')$, leading to a contradiction. To see why, suppose that $U = U''$ and let $f_z^*(U'') = 0$, $e_z^*(U'') = \eta$, and $U_{z,s}^{N*}(U'') = E_z(U_{z,s}^{N*}(U')) \forall z, s$. It is straightforward to check that solution is feasible, satisfies promise keeping and incentive compatibility, where this follows from promise keeping and incentive compatibility of the original solution under U' . From Lemma 8, $J(U_{z,s}^{N*}(U')) = J^{\max}(U_{z,s}^{N*}(U')) \forall z, s$, which means given the weak concavity of $J(\cdot)$ that $J(U_{z,s}^{N*}(U'')) = J^{\max}(U_{z,s}^{N*}(U''))$. Therefore, under the proposed solution,

$$J(U'') = -\pi_a(\eta) \chi + \beta E_z E_s \{ J^{\max}(U_{z,s}^{N*}(U'')) | e = \eta \} \quad \text{and} \quad (\text{A-34})$$

$$U'' = -\eta + \beta E_z E_s \{ U_{z,s}^{N*}(U'') | e = \eta \}. \quad (\text{A-35})$$

One can show by some algebra that $J(U'') = J^{\max}(U'')$ for $U'' > U'$, hence contradicting the fact that U' is the highest continuation value for which $J(U) = J^{\max}(U)$. To show this, use (43) to substitute in for $J^{\max}(U_{z,s}^{N*}(U''))$ in (A-34) and use (A-35) to further substitute in for $U_{z,s}^{N*}(U'')$. Given this contradiction, this establishes that it is necessary that $f_z^*(U') = 0 \forall z$.

Step 2. Given that $f_z^*(U') = 0$, it is necessary that

$$U_{zs}^{*N}(U') = \begin{cases} U' & \text{if } s \leq s' \\ \underline{U}(i^*) & \text{if } s \geq s' \end{cases} \quad (\text{A-36})$$

for some cutoff s' chosen so that (27) binds. To see why, we first establish that (27) binds. Suppose this were not the case. We have already established in Lemma 8 that $U_{zs}^{N*}(U') \leq U' \forall z, s$. Satisfaction of (27) additionally implies that, conditional on z , $U_{zs}^{N*}(U') < U'$ for some s . Given a solution for which (27) does not bind conditional on z , consider an alternative solution which lets $\hat{U}_{zs}^{N*}(U') = \min \left\{ \hat{U}_{zs}^{N*}(U') + \epsilon, U' \right\}$ for some $\epsilon > 0$ arbitrarily small. Such a solution satisfies feasibility and incentive compatibility and makes both the principal and the agent strictly better off, giving the agent some higher continuation value $\hat{U}'$. By analogous arguments as those used in step 1, it is then the case that $J(\hat{U}') = J^{\max}(\hat{U}')$ for $\hat{U}' > U'$, leading to a contradiction. Therefore, (27) must bind.

We now show that the values of $U_{zs}^{N*}(U')$ must satisfy (A-36). Suppose this is not the case for some z . By Lemma 8, $U_{zs}^{N*}(U') \in [\underline{U}(i^*), U']$. Conditional on z , let $\tilde{U} = E_s \{U_{z,s}^N | e_z = \eta\}$. Note that conditional on z , the values of $U_{zs}^{N*}(U)$ must necessarily solve the following program:

$$\max_{\{U_{zs}^N\}_{s \in [0, \bar{s}]}} E_s \{U_{z,s}^N | e_z = \eta\} - E_s \{U_{z,s}^N | e_z = 0\} \quad (\text{A-37})$$

s.t.

$$\tilde{U} = E_s \{U_{z,s}^N | e_z = \eta\} \text{ and} \quad (\text{A-38})$$

$$U_{zs}^N \in [\underline{U}(i^*), U'] . \quad (\text{A-39})$$

This is because if $U_{zs}^{N*}(U')$ does not solve this program, then there exists an alternative value of $\hat{U}_{zs}^{N*}(U')$ which solves this program, is feasible, satisfies promise keeping, satisfies incentive compatibility, and for which (27) is slack. Given the linearity of $J(U)$ for $U \in [\underline{U}((i^*)), \bar{U}((i^*))]$ this solution yields the same welfare to the principal, and this contradicts our previous argument that (27) must bind.

Consider the solution to (A-37) – (A-39). Let λ correspond to the Lagrange multiplier on constraint (A-38). Suppose it were the case that (A-39) does not bind for some s'' and s''' , so that (A-39) does not bind for values of $U_{zs}^{N*}(U')$ which occur with positive probability. First order conditions imply that

$$\frac{\Phi_s(s''', \eta) - \Phi_s(s''', 0)}{\Phi_s(s''', \eta)} = \lambda = \frac{\Phi_s(s'', \eta) - \Phi_s(s'', 0)}{\Phi_s(s'', \eta)}.$$

However, if this is true, then this violates the MLRP property which states that $\Phi_s(s, 0) / \Phi_s(s, \eta)$ is rising in s . This means that (A-39) must binds so that $U_{zs}^{N*}(U')$ either equals $\underline{U}(i^*)$ or U' . This means that from first order conditions,

$$\frac{\Phi_s(s, \eta) - \Phi_s(s, 0)}{\Phi_s(s, \eta)} \geq \lambda \text{ if } U_{zs}^{N*}(U') = U' \text{ and} \quad (\text{A-40})$$

$$\frac{\Phi_s(s, \eta) - \Phi_s(s, 0)}{\Phi_s(s, \eta)} \leq \lambda \text{ if } U_{zs}^{N*}(U') = \underline{U}(i^*). \quad (\text{A-41})$$

Incentive compatibility requires that, conditional on z , $U_{zs}^{N*}(U') = U'$ for some s and $U_{zs}^{N*}(U') = \underline{U}(i^*)$ for some s . Thus, from the MLRP property, there exists some s' such that (A-40) and (A-41) both bind for s' and for which (A-40) and (A-41) imply (A-36).

Step 3. Steps 1 and 2 imply that if there exists a value of $U' > \bar{U}(i^*)$ for which (42) binds, then it must satisfy the following system of equations for some s' :

$$\underline{U}(i^*)(1 - \beta) = -\eta \left(\frac{1/\beta - \Phi(s', 0)}{\Phi(s', \eta) - \Phi(s', 0)} \right), \text{ and} \quad (\text{A-42})$$

$$U'(1 - \beta) = -\eta \left(\frac{1 - \Phi(s', 0)}{\Phi(s', \eta) - \Phi(s', 0)} \right), \quad (\text{A-43})$$

where we have substituted (27) which binds into (24). From Lemma 3 part 4, it follows that (A-42) has at most two solutions. From what we have established in part 1, $s' = \tilde{s}(i^*)$ corresponds to the higher of the two solution since $\tilde{s}(i^*) \geq \hat{s}(\beta)$ for $\hat{s}(\beta)$ defined in Lemma 3. Given (37) and (A-43), such a solution implies that $U' = \bar{U}(i^*)$, contradicting the fact that $U' > \bar{U}(i^*)$. If we instead consider the value of $s' < \hat{s}(\beta)$ which satisfies (A-42), by part 2 of Lemma 3, such a value leads to $U' < \bar{U}(i^*)$, again leading to a contradiction. Therefore, there does not exist $U' \in (\bar{U}(i^*), -\eta/(1 - \beta)]$ for which (42) is an equality. ■

11.2.8 Proof of Lemma 10

Part 1. Suppose by contradiction that for $U \geq \bar{U}(i^*)$ it is the case that $f_z^*(U) = 1$ for some z . It is clear that it must be that $U > \bar{U}(i^*)$, since the arguments in the proof of part 2 of Lemma 9 imply that $f_z^*(\bar{U}(i^*)) = 0 \forall z$. We consider two separate cases in order to rule out that $f_z^*(U) = 1$ for some z . Before doing so, we establish the following preliminary result.

Claim 1. It is the case that $\bar{U}(i^*) \in (-g(0)/(1 - \beta), -\eta/(1 - \beta))$ for $\bar{U}(i^*)$ defined in (38). It is clear from (38) that $\bar{U}(i^*) < -\eta/(1 - \beta)$. To see that $\bar{U}(i^*) > -g(0)/(1 - \beta)$ note that from (38) and (39), $\bar{U}(i^*)$ is decreasing in $\underline{U}(i^*)$ through the implied effect on $\tilde{s}(i^*)$. Since $\underline{U}(i^*) < -g(0)/(1 - \beta)$, it is sufficient to check that $\bar{U}(i^*)$ as defined in (38) exceeds $-g(0)/(1 - \beta)$ if it were the case that the right hand side of (39) were equal to $-g(0)/(1 - \beta)$, which is its maximum. This is guaranteed by the fact that the right hand side of (38) exceeds the right hand side (39).

We can now consider each case.

Case 1. Suppose that $U > \bar{U}(i^*)$, that $f_z^*(U) = 1 \forall z$. Optimality requires that

$$J(U) = E_z \{ -\pi_p \chi - A i_z^*(U) + \beta J(U_z^{F*}(U)) \} = E_z \{ J(-g(i_z^*(U)) + \beta U_z^{F*}(U)) \}. \quad (\text{A-44})$$

Let $\tilde{U} = -g(i_z^*(U)) + \beta U_z^{F*}(U)$ for some z for which $-g(i_z^*(U)) + \beta U_z^{F*}(U) > \bar{U}(i^*)$, where such a z must exist in order to satisfy (24). Let $\tilde{i} = i_z^*(U)$ for the associated value of $i_z^*(U)$. Let $\hat{J}(U|\tilde{i}) = -\pi_p \chi - A \tilde{i} + \beta J((U + g(\tilde{i}))/\beta)$, where it is clear that

$$\hat{J}(\tilde{U}|\tilde{i}) = J(\tilde{U}). \quad (\text{A-45})$$

It is useful to establish the following property of $\hat{J}(\cdot)$. For any two value U' and U'' where $-g(\tilde{i})/(1-\beta) \leq U' < U''$, it must be that

$$\frac{\hat{J}(U''|\tilde{i}) - \hat{J}(U'|\tilde{i})}{U'' - U'} = \frac{J((U'' + g(\tilde{i}))/\beta) - J((U' + g(\tilde{i}))/\beta)}{(U'' - U')/\beta} \quad (\text{A-46})$$

$$\leq \frac{J(U'') - J(U')}{U'' - U'}, \quad (\text{A-47})$$

where we have appealed to the concavity of $J(\cdot)$. We now show that (A-45) cannot hold, which proves by contradiction that it cannot be that $f_z^*(U) = 1 \forall z$.

Case 1a. Suppose that $\tilde{i} \leq i^*$. We have already established that $\bar{U}(i^*) > -g(0)/(1-\beta) \geq -g(\tilde{i})/(1-\beta)$. Moreover, from the proof of part 1 of Lemma 9, $\underline{U}(i^*) \leq -g(i^*) + \beta \bar{U}(i^*)$ since $U_z^{F*}(\underline{U}(i^*)) \leq \bar{U}(i^*)$, and this implies that $\underline{U}(i^*) \leq -g(\tilde{i}) + \beta \bar{U}(i^*) < \bar{U}(i^*)$. Therefore for $\epsilon > 0$ arbitrarily small,

$$J^{\max}(-g(\tilde{i}) + \beta \bar{U}(i^*) + \epsilon) = J(-g(\tilde{i}) + \beta \bar{U}(i^*) + \epsilon) > \hat{J}(-g(\tilde{i}) + \beta \bar{U}(i^*) + \epsilon|\tilde{i}). \quad (\text{A-48})$$

The first equality follows from Lemma 9 since $-g(\tilde{i}) + \beta \bar{U}(i^*) + \epsilon \in [\underline{U}(i^*), \bar{U}(i^*)]$ for arbitrarily small ϵ . The second inequality follows from Lemmas 8 and 9 which imply that if (42) binds, then it cannot be that $U_z^{F*}(U) = \bar{U}(i^*) + \epsilon/\beta > \bar{U}(i^*)$. Letting $U'' = \tilde{U}$ and $U' = -g(i^*) + \beta \bar{U}(i^*) + \epsilon$, substitution of the second inequality in (A-48) into (A-47) implies that (A-45) is violated. Therefore, it is not possible that $\tilde{i} \leq i^*$.

Case 1b. Suppose instead that $\tilde{i} > i^*$. If it is the case that $-g(\tilde{i}) + \beta \bar{U}(i^*) + \epsilon \geq \underline{U}(i^*)$ for arbitrarily small $\epsilon > 0$, then the same arguments as in the case with $\tilde{i} \leq i^*$ imply that this is not possible. If instead $-g(\tilde{i}) + \beta \bar{U}(i^*) < \underline{U}(i^*)$, then

$$J^{\max}(\underline{U}(i^*)) = J(\underline{U}(i^*)) > \hat{J}(\underline{U}(i^*)|\tilde{i}). \quad (\text{A-49})$$

The first equality follows from Lemma 9. The second inequality follows from Lemmas 8 and 9 which imply that if (42) binds, then it cannot be that $i_z^*(U) \neq i^*$. Because $-g(\tilde{i})/(1-\beta) <$

$\underline{U}(i^*) < \tilde{U}$ analogous arguments as in the previous case with $\tilde{i} \leq i^*$ imply that (A-49) substituted into (A-47) implies that (A-45) is violated. Therefore, it is not possible for $\tilde{i} > i^*$. Consequently, it cannot be that $f_z^*(U) = 1 \forall z$.

Case 2. Suppose that $U \geq \bar{U}(i^*)$ and that $f_z^*(U) = 1$ for some z . By analogous arguments as in case 1, optimality requires that

$$J(U) = E_z \left\{ \frac{f_z^*(U) J(-g(i_z^*(U)) + \beta U_z^{F*}(U)) + (1 - f_z^*(U)) J(-e_z^*(U) + \beta E_s \{U_{zs}^{N*}(U) | e_z\})}{f_z^*(U) + (1 - f_z^*(U))} \right\}.$$

Moreover, the concavity of $J(\cdot)$, together with optimality imply that

$$J(U) = E_z f_z^*(U) J(U') + (1 - (E_z f_z^*(U))) J(U'') \quad (\text{A-50})$$

for

$$\begin{aligned} U' &= E_z \{ -g(i_z^*(U)) + \beta U_z^{F*}(U) | f_z^*(U) = 1 \} \text{ and} \\ U'' &= E_z \{ -e_z^*(U) + \beta E_s \{ U_{zs}^{N*}(U) | e_z = e_z^*(U) \} | f_z^*(U) = 0 \} \end{aligned}$$

so that

$$U = E_z f_z^*(U) U' + (1 - (E_z f_z^*(U))) U''.$$

Note that it is necessary given the discussion of case 1 that $U' < \bar{U}(i^*)$, which implies given that $U > \bar{U}(i^*)$ that $U'' > \bar{U}(i^*)$. However, Lemma 9 implies that $J(U)$ is strictly concave between U' and U'' since $U' < \bar{U}(i^*) < U''$. This means that $J(U)$ must exceed the right hand side of (A-50). Therefore, it is not possible that $f_z^*(U) = 1$ for some z .

Part 2. Suppose that $U \geq \bar{U}(i^*)$ and $U_{zs}^{N*}(U) \leq \bar{U}((i^*))$ for some s, z . To prove that (45) holds we consider three different cases.

Case 1. Suppose that for the given z for which $U_{zs}^{N*}(U) \leq \bar{U}((i^*))$ for some s , it is the case that $e_z^*(U) = 0$. We can prove that in this situation if $U_{zs}^{N*}(U) \leq \bar{U}((i^*))$ for some s , then $U_{zs}^{N*}(U) = \bar{U}((i^*))$. Let us first establish that if $e_z^*(U) = 0$ then U is on a portion of $J(\cdot)$ for which $J(U - \epsilon) \geq J(U)$ for $\epsilon > 0$ arbitrarily small, so that it is not possible to make both the principal and the agent strictly better off. Suppose this were not the case and that $J(U - \epsilon) < J(U)$ so that $J(\cdot)$ is upward sloping. Suppose we instead chose $f_z = 1$ and $i_z = 0$ whenever $e_z = 0$, while preserving the rest of the solution. This alternative makes the agent strictly worse off and the principal strictly better off while satisfying feasibility and all incentive compatibility constraints. Therefore, it is not possible that $J(U - \epsilon) < J(U)$. Now suppose that $U_{zs}^{N*}(U) < \bar{U}((i^*))$ for some s . Then it is possible to increase $U_{zs}^{N*}(U)$ by some $\epsilon > 0$ arbitrarily small and given Lemma 9 this makes the principal and the agent strictly better off. Moreover, it continues to satisfy all feasibility and incentive compatibility constraints. Therefore, it is not for $U_{zs}^{N*}(U) < \bar{U}((i^*))$ for any s in this case.

Case 2. Suppose that for the given z for which $U_{zs}^{N*}(U) \leq \bar{U}((i^*))$, it is the case that

$e_z^*(U) = \eta$ and that (27) does not bind. We show that this case it must be that $U_{zs}^{N*}(U) = \overline{U}((i^*))$. Note that optimality requires that given z ,

$$-\pi_a(\eta)\chi + \beta E_s \{J(U_{zs}^{N*}(U)) | e_z = \eta\} = J(-\eta + \beta E_s \{U_{zs}^{N*}(U) | e_z = \eta\}).$$

Moreover, it is necessary that

$$-\eta + \beta E_s \{U_{zs}^{N*}(U) | e_z = \eta\} > \overline{U}((i^*)), \quad (\text{A-51})$$

since if this is not the true then analogous arguments to those of case 2 of part 1 imply that the solution is suboptimal. Now suppose that (27) does not bind, and consider a perturbation which reduce $U_{zs}^{N*}(U)$ by $\epsilon > 0$ arbitrarily small if $U_{zs}^{N*}(U)$ is above the median value of $U_{zs}^{N*}(U)$ (given z) and which increases $U_{zs}^{N*}(U)$ by $\epsilon > 0$ arbitrarily small if $U_{zs}^{N*}(U)$ is below the median value of $U_{zs}^{N*}(U)$. This perturbation satisfies feasibility, promise keeping, and incentive compatibility, and makes the principal weakly better off given the weak concavity of $J(\cdot)$. If it were the case that it did not make him strictly better off, then this would imply that $J(\cdot)$ is a line connecting all points $U_{zs}^{N*}(U)$ conditional on z . Suppose it were that $U_{zs}^{N*}(U) < \overline{U}((i^*))$ for some s . Then this would necessarily imply given Lemmas 7 and 9 that $U_{zs}^{N*}(U) \leq \overline{U}((i^*)) \forall s$ conditional on z . However, since conditional on z , $f_z^*(U) = 0$, $e_z^*(U) = \eta$, and $U_{zs}^{N*}(U) \leq \overline{U}((i^*)) \forall s$, then this implies that

$$-\eta + \beta E_s \{U_{zs}^{N*}(U) | e_z = \eta\} \leq \overline{U}((i^*)) \quad (\text{A-52})$$

which violates (A-51). To see why (A-52) must hold, note that if conditional on z , $f_z^*(U) = 0$, $e_z^*(U) = \eta$, and $U_{zs}^{N*}(U) \leq \overline{U}((i^*)) \forall s$, then this means that

$$J(-\eta + \beta E_s \{U_{zs}^{N*}(U) | e_z = \eta\}) = J^{\max}(-\eta + \beta E_s \{U_{zs}^{N*}(U) | e_z = \eta\}), \quad (\text{A-53})$$

where (A-53) follows by analogous arguments to those of the proof of part 3 Lemma 8. However, if (A-53) holds, then Lemma 9 implies (A-52).

Case 3. Suppose that for the given z for which $U_{zs}^{N*}(U) \leq \overline{U}((i^*))$, it is the case that $e_z^*(U) = \eta$ and that (27) binds. We show that it is not possible that $U_{zs}^{N*}(U) \in (\underline{U}((i^*)), \overline{U}((i^*)))$. Suppose that conditional on z , $U_{zs}^{N*}(U) \in (\underline{U}((i^*)), \overline{U}((i^*)))$ for some s , and denote by $\varphi \in S$ the values of s for which this is the case. Given z , let $\tilde{U} = E_s \{U_{zs}^N | e_z = \eta, s \in \varphi\}$. Note that conditional on z , these values of $U_{zs}^{N*}(U) \in (\underline{U}((i^*)), \overline{U}((i^*)))$ must necessarily solve the

following program:

$$\max_{\{U_{zs}^N\}_{s \in \varphi}} E_s \{U_{z,s}^N | e_z = \eta, s \in \varphi\} - E_s \{U_{z,s}^N | e_z = 0, s \in \varphi\} \quad (\text{A-54})$$

s.t.

$$\tilde{U} = E_s \{U_{z,s}^N | e_z = \eta, (z, s) \in \varphi\} \quad (\text{A-55})$$

$$U_{zs}^N \in [\underline{U}((i^*)), \overline{U}((i^*))] \quad (\text{A-56})$$

This is because if $U_{zs}^{N*}(U)$ for $s \in \varphi$ does not solve this program, then there exists an alternative value of $\hat{U}_{zs}^{N*}(U)$ which solves this program, is feasible, satisfies promise keeping, satisfies incentive constraints on the principal, and for which (27) is slack. Given the linearity of $J(U)$ for $U \in [\underline{U}((i^*)), \overline{U}((i^*))]$ this alternative yields the same welfare to the principal, and this contradicts our previous argument that (27) must bind.

Consider the solution to (A-54) – (A-56). Let λ correspond to the Lagrange multiplier on constraint (A-38). Suppose it were the case that (A-56) does not bind for some s'' and s''' , so that (A-56) does not bind for values of $U_{zs}^{N*}(U)$ which occur with positive probability. First order conditions imply that

$$\frac{\Phi_s(s'', \eta) - \Phi_s(s'', 0)}{\Phi_s(s'', \eta)} = \lambda = \frac{\Phi_s(s''', \eta) - \Phi_s(s''', 0)}{\Phi_s(s''', \eta)}.$$

However, if this is true, then this violates the MLRP property which states that $\Phi_s(s, 0) / \Phi_s(s, \eta)$ is rising in s . ■

11.2.9 Proof of Lemma 11

We prove the uniqueness of the solution conditional on $U = \underline{U}(i^*)$ and $U = \overline{U}(i^*)$. Suppose first that $U = \overline{U}(i^*)$. The arguments in the proof of part 2 of Lemma 9 show that the solution in this case is uniquely determined and coincides with the one described. Now suppose that $U = \underline{U}(i^*)$. Conditional on $f_z^*(\underline{U}(i^*)) = 1 \forall z$, $U_z^{F*}(\underline{U}(i^*))$ is uniquely determined so as to satisfy the promise keeping constraint. Since

$$U_z^{F*}(\underline{U}(i^*)) = (1 - \tilde{d}(i^*)) \overline{U}(i^*) + \tilde{d}(i^*) \underline{U}(i^*),$$

this means that $\tilde{d}(i^*)$ is uniquely determined.

We are left to consider the possibility that $f_z^*(\underline{U}(i^*)) = 0$ for some z . In this case, Lemma 8 implies that $e_z^*(\underline{U}(i^*)) = \eta$ for such z . Satisfaction of (5) given (7) implies that conditional on some z for which $f_z^*(\underline{U}(i^*)) = 0$, it must be that

$$-\eta + E_s \{U_{z,s}^{N*}(\underline{U}(i^*)) | e_z = \eta\} \geq \beta \underline{U}(i^*) > \underline{U}(i^*),$$

where we have used the fact that from (37) it is the case that $\underline{U}(i^*) < 0$. This means that conditional on $f_z^*(\underline{U}(i^*)) = 0$, the agent receives a continuation value which exceeds $\underline{U}(i^*)$. Satisfaction of (2) thus implies that conditional on some z for which $f_z^*(\underline{U}(i^*)) = 1$, it must be that

$$-g(i^*) + \beta U_z^{F*}(\underline{U}(i^*)) < \underline{U}(i^*). \quad (\text{A-57})$$

However, for such a given z , (A-57) given the definition of $\underline{U}(i^*)$ in (37) implies that

$$-\pi_p \chi - Ai^* + \beta J(U_z^{F*}(\underline{U}(i^*))) < J(\underline{U}(i^*))$$

which is not incentive compatible for the principal since (3) is violated. Therefore, if $U = \underline{U}(i^*)$, then $f_z^*(\underline{U}(i^*)) = 1 \forall z$ and $U_z^{F*}(\underline{U}(i^*))$ is uniquely determined. ■

11.2.10 Proof of Lemma 12

Equations (56), (59), and (60) implicitly define the function $\bar{J}(i)$. Given the definition established in Section 10.3, the right hand side of (59) corresponds to $-\omega(\beta, \tilde{s}(i))$. Let $\Psi(i)$ and $\zeta(\tilde{s}(i))$ correspond to the right hand sides of (56) and (60), respectively. It is clear by implicit differentiation of (60) with respect to i that

$$(1 - \beta) \bar{J}''(i) = \zeta''(\tilde{s}(i)) \tilde{s}'(i) + \zeta'(\tilde{s}(i)) \tilde{s}''(i). \quad (\text{A-58})$$

If $i < i^*$, $\tilde{s}'(i) > 0$ from the proof of part 3 of Proposition 4 and $\zeta'(\tilde{s}(i)) > 0$ from the proof of part 4 of Proposition 4. Therefore, $\bar{J}''(i)$ has the same sign as

$$\frac{\zeta''(\tilde{s}(i))}{\zeta'(\tilde{s}(i))} + \frac{\tilde{s}''(i)}{\tilde{s}'(i)}. \quad (\text{A-59})$$

By some algebra,

$$\frac{\zeta''(\tilde{s}(i))}{\zeta'(\tilde{s}(i))} = \left(\frac{\Phi_{ss}(\tilde{s}(i), \eta)}{\Phi_s(\tilde{s}(i), \eta)} + 2 \frac{\Phi_s(\tilde{s}(i), \eta)}{[1/\beta - \Phi(\tilde{s}(i), \eta)]} \right) < 2 \frac{\Phi_s(\tilde{s}(i), \eta)}{1/\beta - \Phi(\tilde{s}(i), \eta)}, \quad (\text{A-60})$$

where the second inequality follows from the fact that $\Phi_{ss}(\tilde{s}(i), \eta) < 0$.

To characterize $\tilde{s}''(i)/\tilde{s}'(i)$, note that $\tilde{s}(i)$ must satisfy

$$\Psi(i) = -\omega(\beta, \tilde{s}(i)) \quad (\text{A-61})$$

so that by implicit differentiation,

$$\tilde{s}'(i) = -\frac{\Psi'(i)}{\omega_s(\beta, \tilde{s}(i))}. \quad (\text{A-62})$$

Implicitly differentiating (A-61) twice with respect to i yields

$$\tilde{s}''(i) = -\frac{\Psi''(i) + \omega_{ss}(\beta, \tilde{s}(i)) \tilde{s}'(i)}{\omega_s(\beta, \tilde{s}(i))}. \quad (\text{A-63})$$

We can show that $\Psi''(i) > 0$, which implies that $\tilde{s}''(i) < 0$. To determine the sign of $\Psi''(i)$, differentiate the right hand side of (56) twice with respect to i to achieve:

$$\Psi''(i) = \frac{-g''(i)(\pi_p - \pi_a(\eta))\chi - 2A\Psi'(i)}{(\pi_p - \pi_a(\eta))\chi + Ai} > 0 \quad (\text{A-64})$$

where we have used the fact that $\Psi'(i) < 0$ for $i < i^*$ to establish the inequality. Since $\Psi'(i) < 0$ and $\Psi''(i) > 0$, we can combine (A-62) and (A-63), to achieve

$$\frac{\tilde{s}''(i)}{\tilde{s}'(i)} = \frac{\Psi''(i) - \omega_{ss}(\beta, \tilde{s}(i)) \frac{\Psi'(i)}{\omega_s(\beta, \tilde{s}(i))}}{\Psi'(i)} \leq -\frac{\omega_{ss}(\beta, \tilde{s}(i))}{\omega_s(\beta, \tilde{s}(i))}. \quad (\text{A-65})$$

Substituting (A-15) into the right hand side of (A-65) taking into account that the second term on the right hand side of (A-15) is positive implies that

$$\frac{\tilde{s}''(i)}{\tilde{s}'(i)} < -2 \frac{\Phi_s(s, 0) - \Phi_s(s, \eta)}{\Phi(s, \eta) - \Phi(s, 0)}. \quad (\text{A-66})$$

Given that $\tilde{s}(i) \geq \hat{s}(\beta)$ for $\hat{s}(\beta)$ defined in (A-12) it follows that the sum of the right hand sides of (A-60) and (A-66) is negative which implies that (A-59) is negative so that $\bar{J}''(i) < 0$ for $i < i^*$. ■