

Selection-Free Predictions in Global Games with Endogenous Information and Multiple Equilibria*

George-Marios Angeletos

MIT and NBER

Alessandro Pavan

Northwestern University

October 18, 2012

Abstract

Global games with endogenous information often exhibit multiple equilibria. In this paper we show how one can nevertheless identify useful predictions that are robust across all equilibria and that could not have been delivered in the common-knowledge counterparts of these games. Our analysis is conducted within a flexible family of games of regime change, which have been used to model, *inter alia*, speculative currency attacks, debt crises, and political change. The endogeneity of information originates in the signaling role of policy choices. A novel procedure of iterated elimination of non-equilibrium strategies is used to deliver probabilistic predictions that an outside observer—an econometrician—can form under arbitrary equilibrium selections. The sharpness of these predictions improves as the noise gets smaller, but disappears in the complete-information version of the model.

JEL Classification Numbers: C7, D8, E5, E6, F3.

Key Words: global games, multiple equilibria, endogenous information, robust predictions.

*This paper grew out of prior joint work with Christian Hellwig and would not have been written without our earlier collaboration; we are grateful for his contribution during the early stages of the project. Previous versions circulated under the title “Robust Predictions in Global Games with Multiple Equilibria: Defensive Policies Against Currency Attacks” and “Preempting Speculative Attacks: Robust Predictions in a Global Game with Multiple Equilibria.” We are thankful to the editor, Gadi Barlevy, and various anonymous referees for suggestions that helped us improve the paper. The usual disclaimer applies.

1 Introduction

In the last fifteen years, the global-games methodology has been used to study a variety of socio-economic phenomena, including currency attacks, bank runs, debt crises, political change, and party leadership.¹ Most of the appeal of this methodology for applied work comes from the fact that it provides a toolkit to arrive at unique-equilibrium selection in settings with coordination problems and self-fulfilling beliefs. This selection facilitates positive and normative predictions that were impossible to make as long as these settings were ridden with multiple equilibria.

More recent work, however, has questioned the applicability of such unique-equilibrium selection by showing that multiple equilibria may naturally reemerge once one recognizes the endogeneity of the information structure. Such endogeneity can result from economic mechanisms that are often central to the applications under consideration. Examples include the signaling role of policy interventions (e.g., Angeletos, Hellwig and Pavan, 2006); the aggregation of information through prices (Angeletos and Werning, 2006; Hellwig, Mukherji, and Tsyvinski, 2006; Ozdenoren and Yuan, 2008); the manipulation of information through propaganda (Edmond, 2011); and learning in dynamic settings (Angeletos, Hellwig and Pavan, 2007; Chassang 2010).

In the present paper, we show how global-game techniques may continue to deliver tight and useful predictions even when the endogeneity of information precludes equilibrium uniqueness. We illustrate this possibility within a flexible family of games of regime change in which the endogeneity of the information originates from the signaling role of the actions of a policy maker. Applications may include a central bank trying to defend a currency peg against a speculative attack; a government facing a self-fulfilling debt crisis; a party leader trying to preempt a defection by party members and donors; or a dictator seeking to prevent a revolution.

The backbone of our framework is therefore similar to that in Angeletos, Hellwig and Pavan (2006), combining a global coordination game with a signaling game.² The contribution, however, is distinct. That earlier work guessed and verified the existence of a particular set of equilibria, which sufficed for establishing equilibrium multiplicity, but did not identify any predictions that hold true across the entire equilibrium set. By contrast, the present paper develops a novel procedure of iterated elimination of non-equilibrium strategies, which need not help establish the existence of

¹The global-games methodology was pioneered by Carlsson and van Damme (1993) and then extended by Morris and Shin (2003) and Frankel, Morris and Pauzner (2003). Applications include Morris and Shin (1998), Corsetti, Dasgupta, Morris and Shin, (2004), and Guimaraes and Morris (2006) for currency crises; Rochet and Vives (2004) and Goldstein and Pauzner (2005) for bank runs; Corsetti, Guimaraes and Roubini (2006) and Zwart (2007) for debt crises; Chamley (1999) and Dasgupta (2006) for investment dynamics; Edmond (2006) for political change; and Dewan and Myatt (2007) for party leadership.

²To be precise, our framework is more general in that it allows for a more flexible payoff structure. This generalization permits us to nest novel applications and to illustrate the broader applicability of our approach. The key contribution of the present paper, however, is not in this generalization per se, but rather in the nature of the results, as explained in the main text.

any particular equilibria, but nevertheless identifies a set of predictions that must hold true in any candidate equilibrium.

This explains why, not only the spirit of the contribution, but also the theoretical arguments in the present paper are closer to standard global games than to our earlier work. In particular, the procedure we develop in this paper shares two important similarities with the iterated deletion of dominated strategies in standard global games (e.g., Carlsson and van Damme, 1993, Morris and Shin, 2002). First, contagion effects emerge across different states of nature because, and only because, of the incompleteness of information. Second, these contagion effects permit one to iteratively rule out more and more strategy profiles by showing that they are inconsistent with equilibrium reasoning.

At the same time, there is an important difference. In standard global games, the players' information about the state of nature—and the associated belief hierarchies—are exogenous. By contrast, beliefs are endogenous in our setting, because of the signaling role of policy. As a result, our procedure puts iteratively tighter bounds, not only on the strategies that can be played in the coordination game among the agents (the receivers), but also on the strategies that can be played by the policy maker (the sender) and thereby on the endogenous belief hierarchies that obtain in the coordination game. This explains the complexity and the novelty of the procedure.

Once the strategy profiles that survive this procedure have been identified, the next step is to translate the results into probabilistic statements that can guide empirical work. To this goal, we introduce an arbitrary sunspot (correlation device) whose realization determines the equilibrium being played. We then ask what restrictions the theory imposes on the joint distribution of fundamentals (which are henceforth identified with the type of the policy maker), policy choices, and regime outcomes from the perspective of an outside observer—say, an econometrician—who is uncertain about which equilibrium is played (that is, who does not know the realization of the sunspot variable).

We are thus able to reach the following testable predictions that hold irrespective of equilibrium selection (i.e., irrespective of the sunspot distribution).

- (i) The probability of regime change is monotone in the fundamentals: weaker policy types face a higher probability of regime change.
- (ii) The probability of policy interventions is non-monotone in the fundamentals: the policy maker intervenes when his type is neither too strong nor too weak.
- (iii) The “need” for policy intervention vanishes as agents become better informed about the type of the policy maker: for all positive-measure sets of types, the probability of policy intervention vanishes as the precision of the agents' information grows.
- (iv) The possibility of multiple equilibria hinges on the possibility of policy intervention: if the

policy maker could commit to a particular policy before observing his type, then he could also guarantee a unique equilibrium. Nonetheless, the policy maker can prefer such commitment over discretion only in so far he expects his type to be strong: weak types are always better off with the option to intervene, despite the fact that this option introduces multiple equilibria.

A number of remarks are worth making regarding the predictions documented above and the overall contribution of the paper.

Although the aforementioned predictions seem quite intuitive, none of them could have been made on the basis of the complete-information variant of the model: that variant is ridden with so many equilibria that “almost anything goes”. What is more, even though the equilibrium set in our framework is not a singleton, it features a sharp discontinuity that is reminiscent of that in standard global games: in the limit as the noise vanishes, the set of equilibrium outcomes of our framework is a measure-zero subset of its complete-information counterpart. These points underscore how global-game techniques retain a strong and useful selection bite in our framework despite the endogeneity of information and the ensuing equilibrium multiplicity.

The predictions documented above are, of course, also satisfied in Angeletos, Hellwig and Pavan (2006), for the latter is a special case of the framework considered here. This, however, does not mean that one could have reached these predictions from that earlier work. The guess-and-verify arguments used in that earlier work were sufficient for establishing the existence of particular equilibria, but were inadequate to identify the set of necessary conditions that must be satisfied by *any* candidate equilibrium. This explains why the theoretical arguments contained in the present are different from those in that earlier work, and are instead closer to the type of iterative elimination arguments used in standard global games.

Identifying necessary conditions that hold for all equilibria does not directly translate to useful testable predictions. One must first study how these necessary conditions map into restrictions on fictitious data generated by the model. The second step we take in this paper is therefore an integral part of our contribution, even though it is not specific to global games.

The above point also explains why the predictions we deliver are only probabilistic: as players can randomize across multiple equilibria, the predicted relations between exogenous fundamentals and endogenous equilibrium outcomes is stochastic. This property, in turn, offers an interesting contrast to unique-equilibrium models. In such models, the theory typically delivers a *deterministic* relation between exogenous fundamentals and endogenous outcomes. Before taking the theory to the data, the econometrician must therefore add some randomness in the form of a residual.³ This residual is most often justified either as measurement error or as explanatory variables that were left outside the theory. By contrast, our approach permits one to accommodate this residual as an

³For example, think of the theory delivering the prediction $Y = \beta X$, and the econometrician then testing the relation $Y = \beta X + \varepsilon$, where Y are the endogenous variables, X the exogenous fundamentals, and ε the residual.

integral part of the theory, as the empirical counterpart of random equilibrium selection.⁴

Equilibrium poses restrictions, not only on the relation between fundamentals, policy choices and regime outcomes, but also on the relation of these objects to the agents' hierarchy of beliefs. Predictions that have to do with the structure of beliefs may be of special interest to the theorist. Nonetheless, such predictions seem of limited value for applied purposes, since beliefs—and especially higher-order beliefs—are unlikely to be contained in the data that are typically available in applied work. Similarly, predictions that tie the effectiveness of policy interventions to, say, 12th order beliefs are also of little value to real-world policy makers. This explains why our analysis concentrates on predictions about the joint distribution of a limited set of variables—fundamentals, policy choices, and regime outcomes—that seem of interest for practical purposes.

This last point also provides a possible re-interpretation of our results. In games with exogenous information, Weinstein and Yildiz (2007) and Penta (2012) have shown that equilibrium multiplicity is degenerate as long as one considers a sufficiently rich topology over belief hierarchies. It is unclear whether a variant of that result applies to the type of environments with endogenous information that we are interested in. But even if it does, the ultimate question of interest here is not per se the determinacy of the equilibrium, but rather the sharpness of the predictions an outside observer can make on the basis of limited data about the belief hierarchy.

In particular, we cannot rule out the possibility that one can induce unique-equilibrium selection in our setting by considering sufficiently rich perturbations of the information structure, including perturbations of the signaling technology. Nonetheless, to the extent that different equilibria are selected by different perturbations, as it is indeed the case in Weinstein and Yildiz (2007), then the essence of our results will survive: one would only have to re-interpret the uncertainty the econometrician faces about equilibrium selection in our setting with the uncertainty he may face about the relevant perturbation.

Layout. The rest of the paper is organized as follows. Section 2 sets up our framework. Section 3 constructs the aforementioned procedure of iterated deletion of non-equilibrium strategies and shows how this procedure identifies a set of tight necessary conditions for the entire equilibrium set. Section 4 translates these conditions into probabilistic predictions about the relation between fundamentals, policy choices and regime outcomes. Section 5 studies the equilibrium value of the option to intervene. Section 6 contrasts the incomplete-information game to its common-knowledge counterpart. Section 7 concludes. All proofs are in the Appendix.

⁴The spirit of this point is similar to that in Chari and Kehoe (2003). That paper consider a setting with a unique equilibrium and herding dynamics. In that setting, the econometrician could face significant uncertainty about the relevant economic outcomes (the occurrence of a crisis) conditional on the observable fundamentals (the strength of the regime), not because of equilibrium multiplicity, but because he could not observe the precise signals observed by the agents and the precise sequence in which agents move.

2 Model

The economy is populated by a big player, who seeks to influence the fate of a regime, and a continuum of small atomistic players, who must choose whether or not to attack the regime (i.e., to take actions that favor the status quo or that favor regime change). To fix ideas, we think of the big player as a “policy maker,” and refer to the small players as to the “agents”. We index the latter by i , assume that they are distributed uniformly over $[0, 1]$, and denote by $A \in [0, 1]$ the measure of the agents attacking the regime (the aggregate size of the “attack”).

Depending on the application of interest, the policy maker could be a central bank trying to avoid the devaluation of a currency (Obstfeld, 1996; Morris and Shin, 1998), a debtor trying to convince creditors to roll over their loans (Calvo, 1988; Corsetti, Guimaraes and Roubini, 2006, Zwart, 2007); a dictator trying to prevent political unrest (Edmond, 2006); or a party leader trying to keep the party united (Dewan and Myatt, 2007). As we explain below, the payoff structure we assume is sufficiently flexible to permit any of these possible interpretations of our framework.⁵

Fundamentals, policy actions, and regime outcome. The payoff structure is parametrized by an exogenous random variable $\theta \in \mathbb{R}$. This variable may affect the strength of the status quo, the policy maker’s preferences, or the agents’ costs and benefits from regime change. The realization of this variable is known to the policy maker, but is only imperfectly observed by the agents. As in the related literature, hereafter we refer to θ interchangeably as to the underlying “fundamentals” and the policy maker’s “type.”

Before agents move, the policy maker can take a costly action in an attempt to influence the agents’ behavior and the fate of the regime. In the context of speculative currency attacks, think of capital controls and monetary interventions that raise domestic interest rates. In the context of debt crises, think of fiscal austerity measures and structural reforms. In the context of a dictatorial regime, think of measures that strengthen the militia, increase the stakes for the supporters of the regime, suppress dissent, or appease the general public with pro-democratic reforms. Finally, in the context of party leadership, think of various concessions and promises made by a party leader in order to discourage his fellow party members (or donors) from defecting.

We capture the action of the policy maker by a variable $r \in [\underline{r}, +\infty)$, which is under the direct control of the policy maker. The regime outcome need not be under his direct control, but can be influenced by his policy choice. We assume that regime change occurs if $R(\theta, r, A) \leq 0$, and does not occur if $R(\theta, r, A) > 0$, where R is a continuous function, strictly increasing in θ , strictly decreasing in A , and nondecreasing in r . These monotonicities—which mean that the chances the status quo survives increase with the fundamentals, decrease with the size of the attack, and

⁵Our framework is flexible but, of course, too abstract to accommodate the institutional details of any particular application. Adding such details may complicate the analysis. However, it may also tighten the predictions, for specific applications can justify tighter assumptions on the payoff structure.

(weakly) increases with the policy maker's action—are quite natural given the class of applications we have in mind. For example, in the context of currency crises, the above assumptions capture the idea that devaluation is more likely when the size of the speculative attack is larger and less likely when the policy maker succeeds in securing a larger amount of reserves. In the context of political unrest, they capture the idea that the regime's probability of survival decreases with the size of the insurgent group and increase with the amount of resources the regime spends on strengthening its police and militias.⁶

Policy maker's payoff. The policy maker's payoff is given by the function

$$U(\theta, r, A) = \begin{cases} W(\theta, r, A) & \text{if } R(\theta, r, A) > 0 \\ L(\theta, r) & \text{if } R(\theta, r, A) \leq 0. \end{cases}$$

We assume that the functions W and L are continuously differentiable and satisfy the following conditions:⁷

1. $W_r(\theta, r, A) < 0$ and $L_r(\theta, r) < 0$;
2. $W_A(\theta, r, A) \leq 0$;
3. there exists a threshold $\underline{\theta}$, defined below, such that $W(\theta, r, 0) - L(\underline{r}, \theta)$ is strictly increasing in θ for all $\theta \geq \underline{\theta}$ and all $r > \underline{r}$.
4. $W(\theta, r, A) - L(\theta, r) > 0$ if $R(\theta, r, A) > 0$.

The first assumption simply means that policy interventions are costly, both in case the status quo survives and in case of regime change. Depending on the application of interest, this assumption may reflect the distortionary effects of higher domestic interest rates (in the context of currency crises), or the economic and political costs of fiscal austerity (in the context of sovereign debt crises).

The second assumption means that, conditional on the regime surviving, the policy maker (weakly) prefers a smaller attack. The case where W is independent of A corresponds to a situation where the policy maker cares about the size of the attack A only in so far the latter affects the fate of the regime, as in the case of a political candidate whose payoff is only determined by whether or not he wins an election or, more generally, of whether or not he retains power.

The third assumption is an increasing-difference condition akin to the role played by familiar single-crossing conditions in signaling games: if setting the minimal-cost policy $\underline{r}$ leads to regime change while raising the policy to some level $r > \underline{r}$ spares the status quo from an attack and results

⁶Note that the regime outcome is a deterministic function of (θ, r, A) , implying that in equilibrium the policy maker faces no uncertainty about regime outcomes. We expect, however, our results to extend to a setting where, with a small probability, the regime collapses for exogenous reasons—e.g., think of devaluation occurring as a result of a large attack by noise traders.

⁷Throughout, for any function $f : \mathbb{R}^n \rightarrow \mathbb{R}$ we let f_x denote its partial derivative with respect to argument x .

in the status quo being preserved, higher types have a stronger incentive to raise the policy than lower types.

The fourth assumption means that the policy maker would not prefer to see the regime collapse when it survives. This assumption is trivially satisfied when the fate of the regime is directly controlled by the policy maker, as typically assumed in models of currency attacks.⁸ To capture a richer set of applications, such as political change, we are allowing the regime outcome to be beyond the direct control of the policy maker; to get sharper predictions, we are ruling out the possibility that the policy maker's preferences for the fate of the regime are negatively correlated with the regime's strength.

Finally, note that L is not allowed to depend on A : conditional on regime change, the policy maker's payoff is assumed to be independent of the size of the attack. This assumption is more restrictive than the preceding ones, but need not be unnatural. For example, once a dictator has been thrown out power, he might not care anymore about the exact size of the revolt that overthrew him—in fact, he might actually be dead! In the context of currency crises, on the other hand, this assumption may be motivated by the idea that the cost of defending the peg against a speculative attack are largely waived if the central bank decides not to defend. Putting aside these possible justifications, the reason we impose this assumption is that it facilitates an important step in our characterization procedure: it guarantees that, along any equilibrium, policy interventions signal the policy maker's confidence that regime change will not occur in equilibrium.⁹

The agents' payoff. As standard in coordination games, what matters for the agents is the payoff differential between attacking and not attacking. Consequently, we normalize the payoff from not attacking to zero and let the payoff that each agent obtains in case he attacks be given by

$$u(\theta, r, A) = \begin{cases} -Q(r) & \text{if } R(\theta, r, A) > 0 \\ Z(\theta, r) - Q(r) & \text{if } R(\theta, r, A) \leq 0. \end{cases}$$

The term $Q(r) > 0$ captures the various costs of attacking in case the attack fails and the regime is maintained, while the term $Z(\theta, r) > 0$ captures the gross benefit of attacking when the attack is successful.¹⁰ The functions Z and Q are bounded, continuously differentiable, and satisfy $Z_\theta, Z_r \leq 0 \leq Q_r$, for all (θ, r) . In the context of a currency attack, for example, $Q(r)$ represents the transaction and other costs that a speculator has to bear in order to attack, while $Z(\theta, r)$ represents

⁸To capture the possibility that the fate of the regime is directly controlled by the policy maker, we can add a final stage during which the policy maker decides whether or not to abandon the status quo after observing the size of attack A . The policy maker would then decide to abandon the status quo if and only if it is in his interest to do so; that is, $R(\theta, r, A) > 0$ if and only if $W(\theta, r, A) > L(\theta, r)$.

⁹This assumption can thus readily be replaced with the weaker assumption that, conditional on regime change, the policy maker's optimal choice is $r = \underline{r}$.

¹⁰Our results extend to the case where Z also depends on the size of the attack A , to the extent that this dependence preserves strategic complementarity in the agents' actions. However, because the exposition is heavier in that case, we restrict attention here to the case where Z depends only on (θ, r) .

the devaluation premium; the dependence of Q on r then captures the ability of the government to manipulate these costs, while the dependence of Z on θ and r captures the possibility that the shadow value of the peg increases with either the quality of the fundamentals or various fiscal reforms that can be undertaken by the government.

Timing and information. The game evolves through two phases. In the first phase, the policy maker sets the policy r after learning θ . In the second phase, agents decide simultaneously whether or not to attack, after observing the policy r , and after receiving private signals $x_i = \theta + \sigma \xi_i$ about θ ; the scalar $\sigma \in (0, \infty)$ parameterizes the quality of the agents' information, while ξ_i is noise, i.i.d. across agents and independent of θ , with a continuous probability density function ψ strictly positive and differentiable over the entire real line, with corresponding cumulative distribution function Ψ . The common prior about θ is uniform over the entire real line.¹¹

Dominance regions. As in all global games, the selection power of incomplete information hinges on the introduction of dominance regions. The ones we assume here are based on two intuitive properties. When the fundamentals are sufficiently weak (low θ), regime change is inevitable, irrespective of the size of the attack and of the level of policy intervention. Likewise, when the fundamentals are sufficiently strong (high θ), the status quo survives irrespective of the size of the attack and of policy interventions. In either case, the policy maker refrains from intervening. These properties are embedded in our model as follows.

First, we assume that there exist finite thresholds $\underline{\theta}$ and $\bar{\theta}$, with $\underline{\theta} < \bar{\theta}$, such that $R(\underline{\theta}, r, 0) = 0 = R(\bar{\theta}, r, 1)$ for any r . This assumption identifies the interval $(\underline{\theta}, \bar{\theta}]$ with the “critical region” of fundamentals over which the regime outcome hinges on the size of the attack,¹² and guarantees that it is dominant for the policy maker to set $r = \underline{r}$ whenever $\theta \leq \underline{\theta}$ (in which case regime change is inevitable). Next, we assume that $\lim_{\theta \rightarrow +\infty} \rho(\theta) = \underline{r}$, where $\rho(\theta)$ is defined as the maximal level of r that is not strictly dominated by $\underline{r}$ for a policy maker of type θ .¹³ This assumption guarantees that any policy action $r > \underline{r}$ is dominated by inaction ($r = \underline{r}$) for sufficiently high fundamentals.¹⁴

¹¹As in the rest of the literature, this improper prior is used only for convenience; the results generalize to any bounded smooth prior as long as σ , the noise in the agents' private signals, is small enough.

¹²The assumption that this region is invariant to r is only for expositional simplicity.

¹³To understand whether raising the policy to some level $r > \underline{r}$ is strictly dominated by leaving the policy at $\underline{r}$, recall that $W(\theta, r, A)$ is non-increasing in A and that $L(\theta, r)$ is independent of A . It follows that the best-case scenario following the former choice is that no agent attacks ($A = 0$), while the worst-case scenario following the latter choice is that all agents attack ($A = 1$). Furthermore, because regime change is inevitable for $\theta \leq \underline{\theta}$ and because the payoff $L(\theta, r)$ in case of regime change is independent of A and decreasing in r , it follows that $\rho(\theta) = \underline{r}$ for all $\theta \leq \underline{\theta}$. In contrast, for any $\theta \in (\underline{\theta}, \bar{\theta}]$, any r , regime change occurs if all agents attack and does not occur if all agents refrain from attacking, implying that $\rho(\theta) = \sup\{r \geq \underline{r} : W(\theta, r, 0) \geq L(\theta, r)\}$. Lastly, for any $\theta > \bar{\theta}$, regime change never occurs, irrespective of the size of the attack, which means that $\rho(\theta) = \sup\{r \geq \underline{r} : W(\theta, r, 0) \geq W(\theta, \underline{r}, 1)\}$.

¹⁴The role of this condition is to rule out equilibria where interventions occur also for arbitrarily high types. These equilibria seem unrealistic and are not robust, for example, to the possibility that the noise in the agents' signals has bounded support. Instead of invoking the assumption of bounded noise, in the sequel we assume directly that policy intervention is dominated for sufficiently high types.

Finally, we assume that $Z(\theta, r) > Q(\theta, r)$ for all (θ, r) such that $\theta \leq \bar{\theta}$ and $r \leq \rho(\theta)$. This assumption simply means that an agent who expects regime change finds it optimal to attack at least in so far the policy maker does not play a dominated action.

3 Equilibrium

Our equilibrium concept is Perfect Bayesian Equilibrium. Let $r(\theta)$ denote the policy chosen by type θ , $\mu(\theta|x, r)$ the cumulative distribution function of the posterior belief of an agent who receives a signal x and observes a policy r , $a(x, r)$ the action of that agent, $A(\theta, r)$ the corresponding aggregate size of attack, and $D(\theta, r)$ the resulting regime outcome, with $D(\theta, r) = 0$ if $R(\theta, r, A(\theta, r)) > 0$, that is, if the status quo is maintained, and $D(\theta, r) = 1$ if $R(\theta, r, A(\theta, r)) \leq 0$, that is, if regime change occurs.¹⁵ The equilibrium definition can then be stated as follows.

Definition. *An equilibrium consists of a strategy for the policy maker, $r : \mathbb{R} \rightarrow [\underline{r}, +\infty)$, a posterior belief for the agents, $\mu : \mathbb{R}^2 \times [\underline{r}, +\infty) \rightarrow [0, 1]$, a strategy for the agents, $a : \mathbb{R} \times [\underline{r}, +\infty) \rightarrow \{0, 1\}$, and an aggregate attack function, $A : \mathbb{R} \times [\underline{r}, +\infty) \rightarrow [0, 1]$, such that*

$$r(\theta) \in \arg \max_r U(\theta, r, A(\theta, r)) \quad \forall \theta \quad (1)$$

$$\mu(\theta|x, r) \text{ is obtained from Bayes' rule using } r(\cdot) \text{ for any } x \in \mathbb{R} \text{ and any } r \in r(\mathbb{R}) \quad (2)$$

$$a(x, r) \in \arg \max_{a \in \{0, 1\}} a \left[\int_{-\infty}^{+\infty} Z(\theta, r) D(\theta, r) d\mu(\theta|x, r) - Q(r) \right] \quad \forall (x, r) \quad (3)$$

$$A(\theta, r) = \int_{-\infty}^{+\infty} a(x, r) \frac{1}{\sigma} \psi\left(\frac{x-\theta}{\sigma}\right) dx \quad \forall (\theta, r) \quad (4)$$

where $r(\mathbb{R}) \equiv \{r : r = r(\theta), \theta \in \mathbb{R}\}$ is the set of policy interventions that are played in equilibrium. The equilibrium regime outcome is given by

$$D(\theta) \equiv D(\theta, r(\theta)) = \begin{cases} 0 & \text{if } R(\theta, r(\theta), A(\theta, r(\theta))) > 0 \\ 1 & \text{otherwise.} \end{cases}$$

Conditions (1) and (3) require that the policy maker's and the agents' actions be sequentially rational. Condition (4) requires that the aggregate size of attack is the one that obtains by aggregating the strategy of the agents. Finally, condition (2) requires that, on the equilibrium path, the agents' beliefs be pinned down by Bayes' rule.¹⁶

Before proceeding to the characterization of the equilibrium set, we introduce some additional notation. Let $\mathcal{E}(\sigma)$ denote the set of all possible equilibria in the game with quality of information σ . Next, for any $s \geq \underline{r}$, let $\mathcal{E}(s; \sigma)$ denote the set of equilibria in which $r(\theta) \in \{r, s\}$ for all θ ,

¹⁵All equilibrium objects depend on σ , the level of noise, but this dependence is left implicit unless otherwise stated.

¹⁶Note that the equilibrium definition restricts attention to symmetric pure-strategy profiles; as discussed at the end of this section, this is without loss of generality.

meaning that the policy takes either the cost-minimizing value $\underline{r}$ or the value s . For any (θ_1, θ_2) with $\theta_2 \geq \theta_1$, let $X(\theta_1, \theta_2; \sigma)$ be the unique solution to¹⁷

$$\frac{\int_{-\infty}^{\theta_1} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{x-\tilde{\theta}}{\sigma}\right) d\tilde{\theta}}{1 - \Psi\left(\frac{x-\theta_1}{\sigma}\right) + \Psi\left(\frac{x-\theta_2}{\sigma}\right)} = Q(\underline{r}). \quad (5)$$

This threshold identifies the unique signal x at which an agent who believes that $\theta \notin [\theta_1, \theta_2]$ and that regime change occurs if and only if $\theta \leq \theta_1$ is indifferent between attacking and not attacking when observing $r = \underline{r}$. Finally, let

$$B(\theta_1, \theta_2; \sigma) \equiv \Psi\left(\frac{X(\theta_1, \theta_2; \sigma) - \theta_2}{\sigma}\right) \quad (6)$$

denote the aggregate size of attack that obtains when the policy maker's type is θ_2 and agents attack if and only if $x < X(\theta_1, \theta_2; \sigma)$.

3.1 Characterization

We now proceed to state our key two characterization results. The first one uses a novel procedure of iterated deletion of strategy profiles that cannot be part of an equilibrium in order to obtain tighter and tighter bounds on the equilibrium set. These bounds, which are presented in Proposition 1 and formally derived in subsection 3.2 below, play a central role in our analysis: they contain the equilibrium set; they rule out a large set of strategy profiles, including many of those that could have been equilibrium profiles under complete information; and they drive the core predictions we present in Section 4. Yet, these bounds need not always be the sharpest: in general, we cannot rule out the possibility that the equilibrium set is strictly smaller than the set identified by these bounds. We eliminate this ambiguity in Proposition 2 for the case that the policy maker's payoff satisfies a natural single-crossing property.

Proposition 1 (necessary conditions). *The following properties are true for any $\sigma > 0$.*

- (i) *The equilibrium set is given by $\mathcal{E}(\sigma) = \cup_{s \geq \underline{r}} \mathcal{E}(s; \sigma)$.*
- (ii) *If $\mathcal{E}(\underline{r}; \sigma) \neq \emptyset$, then any equilibrium in $\mathcal{E}(\underline{r}; \sigma)$ is such that*

$$D(\theta) = 1 \text{ if and only if } \theta \leq \theta^\#,$$

where $\theta^\# = \theta^\#(\sigma)$ is the unique solution to

$$R\left(\theta^\#, \underline{r}, B(\theta^\#, \theta^\#; \sigma)\right) = 0. \quad (7)$$

¹⁷The fact that equation (5) admits a unique solution is established in Appendix B: see the proof of Property 3 inside the proof of Proposition 9.

(iii) For any $s > \underline{r}$, if $\mathcal{E}(s; \sigma) \neq \emptyset$, then there exists a pair of thresholds (θ_s^*, θ_s'') with $\theta_s'' \geq \theta_s^*$ such that¹⁸

$$\theta_s^* = \inf\{\theta \geq \underline{r} : W(\theta, s, 0) \geq L(\theta, \underline{r})\} \quad (8)$$

and either

$$W(\theta_s'', s, 0) = W(\theta_s'', \underline{r}, B(\theta_s^*, \theta_s''); \sigma) \quad \text{and} \quad R(\theta_s'', \underline{r}, B(\theta_s^*, \theta_s''); \sigma) > 0 \quad (9)$$

or

$$R(\theta_s'', \underline{r}, B(\theta_s^*, \theta_s''); \sigma) = 0. \quad (10)$$

(iv) For any $s > \underline{r}$, if $\mathcal{E}(s; \sigma) \neq \emptyset$, then any equilibrium in $\mathcal{E}(s; \sigma)$ is such that

$$r(\theta) = s \quad \text{only if } \theta \in [\theta_s^*, \theta_s^{**}] \quad \text{and} \quad D(\theta) = \begin{cases} 1 & \text{for } \theta < \min\{\theta_s^*, \theta^\#\} \\ 0 & \text{for } \theta > \theta_s^* \end{cases}$$

where

$$\theta_s^{**} = \theta_s^{**}(\sigma) \equiv \sup\{\theta_s'' \geq \theta_s^* : \theta_s'' \text{ satisfies condition (9) or (10)}\}. \quad (11)$$

Furthermore, $\theta_s^* < \theta^\#(\sigma)$ if and only if $s < r^\#$, where $r^\# = r^\#(\sigma)$ is the unique solution to

$$W(\theta^\#(\sigma), r^\#, 0) = L(\theta^\#, \underline{r}). \quad (12)$$

Part (i) establishes that, in any equilibrium, either the policy is left at $\underline{r}$ by all θ , or it is raised to the same level $s > \underline{r}$ by all types who raise the policy above $\underline{r}$. Parts (ii) characterizes the subset of equilibria in which all types pool on $r = \underline{r}$. It identifies a unique threshold $\theta^\#$ such that, in any such equilibrium, regime change occurs if and only if $\theta \leq \theta^\#$.¹⁹ As implied by Condition (7), this threshold is the lowest threshold θ' for which the regime survives when each agent expects no type of the policy maker to raise the policy above $\underline{r}$ and regime change to occur if and only if $\theta < \theta'$, in which case each agent attacks if and only if he receives a signal $x < X(\theta', \theta'; \sigma)$, which implies that the size of the attack at θ' is given by $B(\theta', \theta'; \sigma)$.²⁰

Next, consider the subset of equilibria in which some type raises the policy to s , for some $s > \underline{r}$. Part (iii) identifies necessary conditions for such an equilibrium to exist. Part (iv) in turn establishes that, in any such equilibrium, there exists a pair of thresholds θ_s^* and θ_s^{**} such that (a) the policy is raised to s only if $\theta \in [\theta_s^*, \theta_s^{**}]$, (b) regime change never occurs for $\theta > \theta_s^*$ and (c) regime change

¹⁸When we say that there exists a threshold θ_s^* satisfying (8), we mean that the set $\{\theta \geq \underline{r} : W(\theta, s, 0) \geq L(\theta, \underline{r})\}$ is non-empty.

¹⁹Our discussion in the main text suppresses the dependence of $\theta^\#$, $r^\#$, and θ^{**} on σ to simplify the exposition, although this dependence remains explicit in the statement of the formal results.

²⁰Not surprisingly, this threshold coincides with the one that obtains in a variant of our game that restricts the policy to $r \in \{\underline{r}\}$ for all θ , which is the case typically studied in global games with exogenous information – see e.g., Morris and Shin (1998) for an application to currency crises.

always occurs for $\theta < \min\{\theta_s^*, \theta^\#\}$.²¹ The threshold θ_s^* identifies the lowest type of the policy maker who prefers raising the policy to s to leaving the policy at $\underline{r}$, when the former choice discourages each agent from attacking and spares the policy maker from regime change, whereas the latter choice leads to regime change. As evident from Conditions (9) and (10), θ_s^{**} in turn identifies the highest type $\theta_s'' \geq \theta_s^*$ who finds it optimal to raise the policy to s when leaving the policy at $\underline{r}$ leads to an attack of size $A = B(\theta_s^*, \theta_s''; \sigma)$. Recall that $B(\theta_s^*, \theta_s''; \sigma)$ is the size of the attack that obtains at $\theta = \theta_s''$ when, after observing $\underline{r}$, each agent attacks if and only if he receives a signal $x \leq X(\theta_s^*, \theta_s''; \sigma)$. In the proof below, we will show that, when θ_s'' is the highest type to raise the policy to s , then it is iteratively dominated for each agent to attack for any $x > X(\theta_s^*, \theta_s''; \sigma)$, which implies that $B(\theta_s^*, \theta_s''; \sigma)$ is the largest attack that type θ_s'' can expect to face in case he decides to leave the policy at $\underline{r}$. It follows that, for type θ_s'' to prefer raising the policy to s than leaving the policy at $\underline{r}$, it must be that the largest attack $B(\theta_s^*, \theta_s''; \sigma)$ that type θ_s'' can possibly expect in case he leaves the policy at $\underline{r}$ either leads to regime change or to a payoff $W(\theta_s'', \underline{r}, B(\theta_s^*, \theta_s''; \sigma))$ that is lower than the payoff $W(\theta_s'', s, 0)$ that type θ_s'' could obtain by raising the policy to s and guaranteeing that no agent attacks.

Note that, while Proposition 1 identifies properties that *any* equilibrium must satisfy, it does not guarantee existence of equilibria satisfying such properties. This incompleteness, however, disappears once we impose the following single-crossing restriction on the policy maker's payoff. Let $\Delta W(\theta; r, A) \equiv W(\theta, r, 0) - W(\theta, \underline{r}, A)$ denote the payoff differential between raising the policy to $r > \underline{r}$ and facing no attack and leaving the policy at $\underline{r}$ and suffering an attack of size A , when the policy maker's type is θ and both choices lead to no regime change.

Single-Crossing Condition (SCC). *The differential $\Delta W(\theta; r, A)$ changes sign at most once as θ increases from $\underline{\theta}$ to $+\infty$: for any $A > 0$ and $r > \underline{r}$, either $\Delta W(\theta; r, A) < 0$ for all $\theta \geq \underline{\theta}$, or there exists a $\theta^+(r, A) > \underline{\theta}$ such that $\Delta W(\theta; r, A) > 0$ if $\theta < \theta^+(r, A)$ and $\Delta W(\theta; r, A) < 0$ for $\theta > \theta^+(r, A)$.*

The SCC thus requires that, in the absence of regime change, the net benefit of reducing the size of the attack from A to zero by raising the policy from $\underline{r}$ to r changes sign only once. Obviously, this condition is trivially satisfied when the policy maker cares about the size of the attack only in so far it affects the regime outcome (i.e., when $W_A(\theta, r, A) = 0$ for all (θ, r, A) , in which case $\Delta W(\theta; r, A) < 0$ for all θ), an assumption that is often made in games of political change that are modeled as a “winner-takes-all election.”)

²¹When $\theta_s^* > \theta^\#$, our procedure of iterated deletion of non-equilibrium strategies does not permit us to pin down the equilibrium regime outcome for $\theta \in [\theta^\#, \theta_s^*]$. What we know is that, if equilibria exist where $\theta_s^* > \theta^\#$, then no type $\theta \in [\theta^\#, \theta_s^*]$ would raise the policy to s . As explained below, this indeterminacy vanishes if one assumes that the policy maker's payoff satisfies the single-crossing condition SCC (see Proposition 2 below), or if one assumes that σ is small (see Proposition 9 in Appendix A).

Proposition 2 (complete characterization). *Suppose SCC holds. Then, the following properties are true for any $\sigma > 0$:*

- (i) $\mathcal{E}(s; \sigma) \neq \emptyset$ if and only if $s \leq r^\#(\sigma)$, where $r^\#(\sigma)$ is the threshold defined in Proposition 1.
- (ii) For any $s \in (\underline{r}, r^\#(\sigma)]$, there exists an equilibrium in $\mathcal{E}(s; \sigma)$ in which $r(\theta) = s$ for all $\theta \in (\theta_s^*, \theta_s^{**}(\sigma)]$ (for all $\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]$ if $\theta_s^* > \underline{\theta}$).

These results complement Proposition 1 in three ways. First, they rule out the possibility of equilibria in which the policy is raised to $s > r^\#$. Second, for any $s \leq r^\#$, they guarantee existence of equilibria in which some type of the policy maker raises the policy to s . Finally, they establish that the thresholds θ_s^* and θ_s^{**} are the sharpest bounds for the set of types who possibly raise the policy to s , in the sense that there always exists an equilibrium in which the policy is raised to s for the entire interval $(\theta_s^*, \theta_s^{**}]$.²²

As the name indicates, the role played by SCC in these results is to guarantee that the policy maker's preferences (and hence his incentives) satisfy a natural ordering condition; when leaving the policy at $\underline{r}$ leads to an attack that is non-increasing in θ , then if type θ_s^{**} prefers raising the policy to s than leaving the policy at $\underline{r}$, so does any type $\theta \in (\theta_s^*, \theta_s^{**}]$.

Remark 1. Proposition 2 can be viewed as a generalization of the results in Propositions 1, 2, and 5 of Angeletos, Hellwig and Pavan (2006): that paper considered a more restrictive payoff specification, which happened to satisfy *SCC*, and established existence of a *particular* set of equilibria in which the policy maker intervenes for intermediate types. None of the results in that paper, however, permits one to identify properties that hold true across *all* equilibria. Apart from the more flexible payoff specification, the key contribution here is thus not Proposition 2 but Proposition 1. It is this novel result, and only this one, that permits one to identify predictions that are robust across the entire equilibrium set, whatever that might be.

Remark 2. In Section 4 below, we show how the properties identified above can be translated into stochastic predictions about policy interventions and regime outcomes as well as into predictions for how the bounds on these variables change with the fundamentals and with the quality of the agents' information. These results hinge on parts (iii)-(iv) of Proposition 1, but depend on Proposition 2 only in so far the results in Proposition 2 rule out the possibility of equilibria in which $s > r^\#$. This possibility is problematic for our purposes because Proposition 1 guarantees monotonicity of the regime outcome for $s \leq r^\#$ but not for $s > r^\#$. The role of *SCC* is precisely to rule out this possibility for any $\sigma > 0$. When *SCC* fails to hold, it is unclear whether equilibria with $s > r^\#$ exist. Even if such equilibria exist, they would not interfere with our predictions as long as

²²The result in Proposition 2 leaves open the possibility that there also exist equilibria in which the policy is raised only for a strict subset of $(\theta_s^*, \theta_s^{**}]$. These equilibria are sustained by the agents following a strategy, after observing the policy maker's choice of not intervening, that is non-monotone in their signals. Although this possibility does not interfere with the predictions we deliver in the sequel, it can be ruled out if one assumes that the noise distribution is log-concave, an assumption which is often made in applications—see Proposition 10 in Appendix A.

they maintain monotonicity of regime outcomes. Furthermore, one can show that, irrespective of whether or not *SCC* holds, equilibria with $s > r^\#$ cannot exist in the limit as $\sigma \rightarrow 0$ (see Proposition 9 in Appendix A). We conclude that the predictions we deliver in the sequel are unlikely to hinge on *SCC*.

3.2 Characterization procedure

We now expand on the series of arguments that lead to Proposition 1. As anticipated, this series of arguments deals with the interaction of two kinds of forces. On the one hand, the incompleteness of information, coupled with the existence of dominance regions, induces a series of contagion effects across different types of the policy maker, as well as across different signals for the agents. On the other hand, the signaling role of policy implies that the information, and the belief hierarchy, in the continuation game that follows any particular policy choice is endogenous to equilibrium play. The contagion effects are reminiscent of those in standard global games, but the endogeneity of information invalidates standard global-game arguments. The interaction of these two forces therefore explains the novelty and complexity of some of the arguments that follow.

We start by considering the subset of equilibria in which all types pool on $\underline{r}$. When this is the case, the observation of $\underline{r}$ conveys no information, and the coordination game that follows this observation is essentially a standard global game. The next lemma then follows from the same kind of arguments as in Morris and Shin (1998, 2003), adapted to the more general environment considered here. (Remark: in all lemmas that follow, the variables $\theta^\#$, $r^\#$, θ_s^* , and θ_s^{**} are those defined in Proposition 1.)

Lemma 1. *In any equilibrium in which no type intervenes, $D(\theta) = 1$ if and only if $\theta \leq \theta^\#$.*

The above lemma establishes part (ii) of Proposition 1. The next two lemmas shift attention to equilibria in which some types raise the policy above $\underline{r}$ and help establish part (i) of the proposition.

Lemma 2. *In any equilibrium in which some type intervenes, there exists a single $s > \underline{r}$ such that $r(\theta) = s$ whenever $r(\theta) \neq \underline{r}$. Furthermore, for that s , $A(\theta, s) = 0$ for all θ .*

Lemma 3. *For any $s > \underline{r}$, if $\mathcal{E}(s; \sigma) \neq \emptyset$, then there exists a $\theta \geq \underline{\theta}$ such that $W(\theta, s, 0) \geq L(\theta, \underline{r})$. Furthermore, any equilibrium in $\mathcal{E}(s; \sigma)$ is such that $D(\theta) = 0$ for all $\theta > \theta_s^*$.*

The key insights behind these two lemmas are simple. First, a type finds it optimal to intervene only if he expects that, by doing so, he will reduce the size of the attack and/or reduce the chances of regime change by strengthening the status quo. Second, if a certain type avoids regime change by raising the policy to r , then any higher type must also be avoiding regime change in equilibrium—for any higher type can always “imitate” any lower type and do at least as well.

These properties are intuitive. Clearly, there must be *some* benefit, in the form of a reduction in the size of the attack and/or a reduction in the probability of regime change, to justify the cost of

intervention. Our model captures this benefit in a stark way, by guaranteeing that *no* attack takes place with probability one following any (equilibrium) intervention. This is because we have ruled out any source of aggregate uncertainty in the size of the attack and in the regime outcome beyond the one introduced by the random type of the policy maker. In the absence of such aggregate uncertainty, equilibrium policy interventions necessarily signal that regime change will not occur (for, otherwise, the policy maker would be better off by not intervening). This in turn implies that, in equilibrium, any type who intervenes does so by selecting the least costly policy among those that are conducive to no regime change, as stated in Lemma 2.

One should not, however, misinterpret Lemma 2 as saying that a single level of policy intervention is consistent with equilibrium behavior. Rather, as established in Proposition 2, there is a continuum of policy levels $s > \underline{r}$ that can be played in some equilibrium. Furthermore, the same type θ may be playing different r in different equilibria, just as the same r may be played by different θ . This point underscores that Lemma 2 has very little positive content on its own right, and anticipates the necessity of the additional arguments we make in Section 4.

Lemma 2 nevertheless helps us index the equilibrium set in a convenient way: any equilibrium that does not belong to $\mathcal{E}(\underline{r}; \sigma)$ necessarily belongs to $\mathcal{E}(s; \sigma)$ for some $s > \underline{r}$. Together with Lemma 1, this lemma therefore establish part (i) of Proposition 1.

Finally, Lemma 3 guarantees that the threshold θ_s^* is an upper bound for the set of types for whom regime change occurs, across all equilibria in $\mathcal{E}(s; \sigma)$. This is because any type above this threshold can always save the regime (and then obtain a payoff higher than in case of regime change) by raising the policy to s and then face no attack. Any type $\theta > \theta_s^*$ who, in equilibrium, leaves the policy at $\underline{r}$ must thus necessarily expect a small attack that does not trigger regime change. Clearly, θ_s^* is also a lower bound on the set of types who possibly raise the policy to s : For any type below this threshold, raising the policy to $r = s$ is dominated by leaving the policy at $\underline{r}$.

Moving on, the next lemma uses a contagion argument to establish that the threshold θ_s^{**} defined in Proposition 1 is an upper bound to the set of types who potentially raise the policy to $r = s$.

Lemma 4. *For any $s > \underline{r}$, if $\mathcal{E}(s; \sigma) \neq \emptyset$, then there exists a $\theta_s'' \geq \theta_s^*$ that satisfies either condition (9) or condition (10) in Proposition 1. Furthermore, any equilibrium in $\mathcal{E}(s; \sigma)$ is such that $r(\theta) = s$ only if $\theta \in [\theta_s^*, \theta_s^{**}]$.*

The proof of Lemma 4 encapsulates contagion effects from very high types, for whom raising the policy is dominated, to lower types, who are spared from the need to raise the policy thanks, and only thanks, to the incompleteness and dispersion of information among the agents. In particular, the fact that raising the policy is dominated for sufficiently high types implies that agents, on their part, find it iteratively dominant to not attack for sufficiently high signals, conditional on observing no policy intervention.²³ The dispersion of information then initiates a contagion effect that triggers

²³That raising the policy to $r = s$ is dominated for arbitrarily high types follows from the assumption that

agents not to attack for lower and lower signals and hence spares the policy maker from the need to raise the policy for lower and lower θ . In the limit, this contagion converges to θ_s^{**} , guaranteeing that all types above this threshold (i) are able to avoid regime change without intervening, and (ii) obtain a higher payoff by leaving the policy at $\underline{r}$ and facing a small attack than by raising the policy to s and face no attack.

Underscoring the power of this contagion effect, the limit threshold θ_s^{**} is arbitrarily close to θ_s^* when σ is small enough (See Proposition 9 in Appendix A) meaning that almost all types for whom policy intervention is not dominated succeed to avoid regime change without the need for costly policy interventions. This is despite the fact that the aforementioned contagion effect is initiated with types that can be arbitrarily high.

Lemma 4 used a contagion argument “from above” to identify a necessary condition for existence of equilibria in which the policy is raised to $r = s > \underline{r}$ by some type and identified an upper bound θ_s^{**} for the set of types who possibly raise the policy to $r = s$. The next lemma uses an alternative contagion argument “from below” to establish that regime change necessarily occurs for any $\theta < \min\{\theta_s^*, \theta^\#\}$.

Lemma 5. *Take any $s > \underline{r}$ and suppose that $\mathcal{E}(s; \sigma) \neq \emptyset$. Then any equilibrium in $\mathcal{E}(s; \sigma)$ is such that $D(\theta) = 1$ for any $\theta < \min\{\theta_s^*, \theta^\#\}$. Furthermore, $\theta_s^* > \theta^\#$ if and only if $s > r^\#$.*

As explain in the Appendix, the contagion effect behind the above result is also present because, and only because, of the dispersion of information. In particular, the fact that, for sufficiently low types, raising the policy is dominated along with the fact that, for these types, regime change is inevitable, implies that agents find it iteratively dominant to attack for sufficiently low signals as long as they do not observe policy intervention. The dispersion of information then initiates a contagion effect such that, conditional on seeing no intervention, agents find it iteratively dominant to attack for higher and higher signals, in which case regime change occurs for higher and higher θ . In the limit, this contagion effect guarantees that regime change occurs for all types below $\min\{\theta_s^*, \theta^\#\}$ in any equilibrium in which the policy is raised to s . This last result is obtained by comparing the agents’ incentives to attack after observing $\underline{r}$ with the corresponding incentives when they expect $r(\theta) = \underline{r}$ for all θ . Because the observation of $\underline{r}$ is most informative of regime change when all types who experience regime change set $r = \underline{r}$, while some of the types who are spared

$\lim_{\theta \rightarrow +\infty} [W(\theta, s, 0) - W(\theta, \underline{r}, 1)] < 0$. Without an assumption of this sort, there may exist equilibria in which the policy maker intervenes even for arbitrarily high types. These equilibria are sustained either by the assumption that the cost of intervention vanishes for sufficiently high types, or by the agents threatening to attack no matter how favorable their signal is when the policy maker fails to intervene. We find either property implausible. Also note that equilibria in which the agents attack no matter their signal when $r = \underline{r}$ are not robust to the following perturbation. Pick any $K > \bar{\theta}$ and any $\delta > 0$ and suppose that with probability δ types $\theta > K$ are forced to set $\underline{r}$ and assume that this event is not observed by the agents. The aforementioned equilibria are not robust to this perturbation, no matter how unlikely this event is (i.e. no matter how big K is).

from regime change raise the policy above $\underline{r}$, the size of attack when setting $r = \underline{r}$ is necessarily larger in any of the equilibria in which some types are expected to raise the policy to $r = s$ than in the pooling equilibria where all types are expected to set $r = \underline{r}$. Hence any type $\theta < \theta^\#$ who does not raise the policy to $r = s$ necessarily experiences regime change in equilibrium. Because raising the policy to $r = s$ is dominated for all $\theta < \theta_s^*$, this implies that regime change occurs *for any* $\theta < \min\{\theta_s^*, \theta^\#\}$ in any equilibrium in which the range of the policy is $\{\underline{r}, s\}$.²⁴

The combination of Lemmas 1 through 5 establishes the results in Proposition 1.

Remark 3. As noted before, Lemma 2, which guarantees that no agent attacks following the observation of an equilibrium policy intervention, hinges on the absence of exogenous aggregate uncertainty. Exogenous aggregate uncertainty could originate in shocks to fundamentals and/or unpredictable shifts in the “sentiment” of some irrational agents that occur after the policy maker has set the policy. While the introduction of such additional uncertainty may deliver a smoother relationship between the probability of regime change and the level of policy intervention, we do not expect our results to be unduly sensitive to our choice of abstracting from such additional uncertainty. For example, it is easy to show that all our results are robust to the introduction of an exogenous random event that triggers regime change independently of r , which one could then interpret either as the result of unfavorable changes in the fundamentals or as the impact of a severe attack by “irrational agents”. Clearly, the same remains true if the probability of this “exogenous” event is decreasing in θ .²⁵ While it could be worthwhile to extend the analysis to more general, and more realistic, sources of aggregate uncertainty, this is beyond the scope of the present paper.

Remark 4. The equilibrium definition we have used rules out mixed strategies for either the policy maker or the agents; it also imposes symmetry on the agents’ strategies. However, from the arguments in the proofs of Lemmas 1-5, it should be clear that none of the conditions identified in these lemmas depends on these restrictions. Indeed, the policy maker can find it optimal to randomize over r only for a zero-measure subset of θ ; because this does not have any effect on the agents’ posterior beliefs about policy and regime outcomes, it cannot affect their best-responses. Similarly, for any r , the agents can find it optimal to randomize over their decision to attack, or to play asymmetrically, only for a zero-measure subset of their signal space; because this does not have any effect on the aggregate size of attack, it does not impact the policy maker’s incentives. Propositions 1 and 2 thus identify properties of *all* equilibrium outcomes, including those sustained by mixed-strategy or asymmetric-strategy profiles.

²⁴That $\theta_s^* < \theta^\#$ if and only if $s < r^\#$ follows from the monotonicity of $W(\theta, s, 0) - L(\theta, \underline{r})$ in θ along with $W_r < 0$.

²⁵Obviously, all relevant thresholds must be adjusted to accommodate the probability of this exogenous event. It is also immediate to see that the probability of such event must not be too high, for otherwise each agent may find it optimal to attack irrespective of his beliefs about the policy maker’s strategy and the behavior of the other agents.

4 Predictions about policy and regime outcomes

We now show how the equilibrium properties identified in the previous section can be translated to meaningful predictions that an outside observer, say an “econometrician”, can make without knowing which particular equilibrium is played. To this goal, we consider two complementary approaches. The first approach formalizes the econometrician’s uncertainty about the equilibrium being played by means of a generic distribution over the set of all possible equilibria, and then proceeds to study the implied distribution of equilibrium outcomes. The second approach studies the extrema of all such possible distributions.

4.1 Probabilistic predictions for arbitrary equilibrium selections

Notwithstanding our earlier remark that *SCC* might not be strictly needed, we henceforth impose *SCC*. We then have that, for any $\sigma > 0$, the equilibrium set is given by $\mathcal{E}(\sigma) = \cup_{s \in [\underline{r}, r^\#(\sigma)]} \mathcal{E}(s; \sigma)$, where $\mathcal{E}(s; \sigma)$ denotes the set of equilibria in which the range of the policy is $\{\underline{r}, s\}$. Different equilibria within the set $\mathcal{E}(s; \sigma)$ are associated with different out-of-equilibrium strategies by the agents. They may also differ in the shape of the policy within the interval $[\theta_s^*, \theta_s^{**}]$, but only if the distribution of the noise in the agents’ information is not log-concave (see the result in Proposition 10 in Appendix A). Nevertheless, all equilibria within the set $\mathcal{E}(s; \sigma)$ are characterized by the same bounds θ_s^* and θ_s^{**} on the set of types who possibly raise the policy, as well as by the same regime outcomes.²⁶ Given the type of predictions we are interested in, from the econometrician’s viewpoint, any distribution over outcomes generated by a random selection over the equilibrium set $\mathcal{E}(\sigma)$ can then be replicated by a random variable $\tilde{s}$ with cumulative distribution function F and support $Supp[F] \subseteq [\underline{r}, r^\#]$ such that one of the pooling equilibria is played when $s = \underline{r}$, while one of the semi-separating equilibria with range $\{\underline{r}, s\}$ is played when $s \in (\underline{r}, r^\#]$. For any σ , we then denote by $\mathcal{F}(\sigma)$ the set of the c.d.f.s. with support $Supp[F] \subseteq [\underline{r}, r^\#(\sigma)]$ that describe the possible beliefs that the external observer may have about the equilibrium being played.²⁷ (Note that the reason why this set depends on the quality σ of the agents’ information is that the upper bound $r^\#(\sigma)$ for the set of possible equilibrium policy interventions depends on σ).

²⁶To be precise, the last statement is true only up to a zero-measure set of types: if the type is exactly equal to the threshold θ_s^* , there is both an equilibrium in $\mathcal{E}(s; \sigma)$ for which the status quo survives and one for which regime change occurs. However, this kind of indeterminacy is irrelevant as long as one forms probabilistic statements over positive-measure sets of fundamentals. To simplify the exposition, in the sequel we impose that the status quo survives when $\theta = \theta_s^*$. Relaxing this assumption does not affect the essence of any of the results; it only complicates the exposition by requiring that Propositions 3 and 5 be restated in terms of probabilistic statements conditional on θ belonging to positive-measure sets rather than conditional on θ taking a specific value. Propositions 4 and 6, on the other hand, are not affected at all.

²⁷As a technical restriction, we assume that $\mathcal{F}(\sigma)$ is compact with respect to the metric $d(\cdot)$ defined, for any pair $F_1, F_2 \in \mathcal{F}$, by $d(F_1, F_2) \equiv \sup \{|F_1(A) - F_2(A)| : A \in \Sigma\}$, where Σ is the Borel sigma algebra associated with the interval $[\underline{r}, \rho(\bar{\theta})]$.

Importantly, note that, because θ is the policy maker's *private* information, the equilibrium being played cannot be a function of θ . This also means that the external observer cannot expect the realization of the random variable $\tilde{s}$ with distribution F to depend of θ . On the other hand, the external observer may be able to observe the policy maker's type θ at the time he/she performs her investigations (for example, the underlying economic fundamentals that parametrize the policy maker's type may become observable ex-post).

Now, for any $\sigma > 0$ and any $s \in [\underline{r}, r^\#]$, let $D_s(\theta; \sigma)$ denote the regime outcome in any of the equilibria in $\mathcal{E}(s; \sigma)$, while, for any $s \in (\underline{r}, r^\#]$, let $\Delta_s \equiv \theta_s^{**} - \theta_s^*$ denote the (Lebesgue) measure of the set of types who potentially raise the policy to $r = s$ in any of the equilibria in $\mathcal{E}(s; \sigma)$.²⁸ Recall that there is no equilibrium in which some type outside the interval $[\theta_s^*, \theta_s^{**}]$ raises the policy to s , while there is an equilibrium in which all types in $[\theta_s^*, \theta_s^{**}]$ raise the policy to $r = s$. Given the uniform prior, Δ_s can thus also be read as (a rescaling of a sharp) bound on the probability that the policy is raised to $r = s$ across all equilibria in $\mathcal{E}(s; \sigma)$.

Next, let I_{premise} denote the indicator function assuming value one if *premise* is true and zero otherwise. For any selection (equivalently, for any belief) $F \in \mathcal{F}(\sigma)$, and any $r > \underline{r}$, then let

$$D(\theta; F, \sigma) \equiv \int D_s(\theta; \sigma) dF(s) \quad \text{and} \quad P(r, \theta; F, \sigma) \equiv \int_{s \geq r} I_{\{\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}} dF(s)$$

denote, respectively, the probability that regime change occurs for type θ and the (maximal) probability that type θ raises the policy to or above r , when the selection is F .²⁹ For any $r > \underline{r}$, any $F \in \mathcal{F}(\sigma)$, then let $\Delta(r; F, \sigma)$ denote the expected Lebesgue measure of the set of types who possibly raise the policy to or above r . Once again, $\Delta(r; F, \sigma)$ can be read as a (rescaling of) the maximal probability that the policy is raised to or above r , when the selection is F . Given two selections $F, F' \in \mathcal{F}(\sigma)$, then let $F' \gg F$ if and only if $F'(s) \leq F(s)$ for all s , with strict inequality for $s \in (\underline{r}, r^\#)$ and equality for $s \in \{\underline{r}, r^\#\}$.

The key predictions about the probability of regime change and of policy interventions that the theory delivers to an external observer who is (potentially) uncertain about which equilibrium is played are summarized in the following proposition.³⁰

Proposition 3 (stochastic predictions). *Equilibrium policies and regime outcomes satisfy the following properties.*

(i) **Non-monotonic policy.** *For any $\sigma > 0$, any $r > \underline{r}$ and any $F \in \mathcal{F}(\sigma)$, there exist thresholds $\theta^\circ = \theta^\circ(r; F, \sigma)$ and $\theta^{\circ\circ} = \theta^{\circ\circ}(r; F, \sigma)$, with $\underline{\theta} < \theta^\circ \leq \theta^{\circ\circ}$, such that $P(r, \theta; F, \sigma) > 0$ only if $r \leq r^\#(\sigma)$ and $\theta \in [\theta^\circ, \theta^{\circ\circ}]$.*

²⁸Recall that θ_s^* is independent of σ .

²⁹Recall that any equilibrium in $\mathcal{E}(s; \sigma)$ is such that $r(\theta) = s$ only if $\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]$ so that $P(r, \theta; F, \sigma)$ is an upper bound on the probability that type θ raises the policy to or above r .

³⁰The predictions that the theory delivers to an observer who knows which equilibrium is played can be read by looking at the special case where the distribution F has a measure-1 mass point at a particular $s \in [\underline{r}, r^\#(\sigma)]$.

(ii) **Monotonic regime outcome.** For any $\sigma > 0$ and any $F \in \mathcal{F}(\sigma)$, $D(\theta; F, \sigma)$ is non-increasing in θ , with $D(\theta; F, \sigma) = 1$ for $\theta \leq \underline{\theta}$ and $D(\theta; F, \sigma) = 0$ for $\theta > \theta^\#(\sigma)$.

(iii) **Impact of “aggressiveness”.** For any $\sigma > 0$ and any $F, F' \in \mathcal{F}(\sigma)$, $F' \gg F$ implies $D(\theta; F', \sigma) > D(\theta; F, \sigma)$ for any $\theta \in (\underline{\theta}, \theta^\#(\sigma))$ such that $D(\theta; F, \sigma) < 1$. Moreover, if $F, F' \in \mathcal{F}(\sigma)$ are such that $F'(r) = F(r)$ and $F'(s) < F(s)$ for all $s \in (r, r^\#(\sigma))$, then $\Delta(r; F', \sigma) \leq \Delta(r; F, \sigma)$.

(iv) **Impact of noise.** Take any F such that $\text{Supp}[F] \subset (\underline{r}, \lim_{\sigma \rightarrow 0+} r^\#(\sigma))$. For any $r > \underline{r}$, $\lim_{\sigma \rightarrow 0+} \Delta(r; F, \sigma) = 0$ whereas, for any θ , any $\sigma, \sigma' > 0$ such that $F \in \mathcal{F}(\sigma) \cap \mathcal{F}(\sigma')$, then $D(\theta; F, \sigma) = D(\theta; F, \sigma')$. Finally, for any $\sigma, \sigma' > 0$ such that $F \in \mathcal{F}(\sigma) \cap \mathcal{F}(\sigma')$, $\sigma' > \sigma > 0$ implies $\Delta(r; F, \sigma') \geq \Delta(r; F, \sigma)$ for all $r \in (\underline{r}, \min \{r^\#(\sigma), r^\#(\sigma')\})$.

Parts (i) and (ii) say that, for any given quality of information $\sigma > 0$ and any selection $F \in \mathcal{F}(\sigma)$, the probability of observing the policy maker raising the policy above r is positive only if $r \leq r^\#(\sigma)$ and θ is intermediate, whereas the probability of regime change is non-increasing in θ and equal to zero for all $\theta > \theta^\#(\sigma)$. These results follow directly from the fact that these properties hold in any of the equilibria of the game (as established in Proposition 1) and hence are preserved in expectation.

Property (iii), on the other hand, can be interpreted as the impact of the “aggressiveness of market expectations”: the higher the level of policy intervention s at which the agents switch to lenient behavior (i.e., refrain from attacking), the higher the cost of policy intervention necessary to prevent an attack, and the smaller the set of types who find it optimal to intervene in equilibrium. Of course, since the level of “aggressiveness” (equivalently, the distribution F) is not necessarily observed by the econometrician, this prediction is hard to test. Yet, it could help the econometrician identify, or estimate, the underlying equilibrium selection, F , from observed data.

Part (iv) says that, holding constant the econometrician’s beliefs F about which equilibrium is played, an increase in the quality of the agents’ information need not have any effect on the probability of regime change, whereas it reduces the probability of observing a policy above r , for any $r > \underline{r}$, with such a probability vanishing in the limit, when information becomes infinitely precise. This result follows from the fact that, for any $s > \underline{r}$, the threshold θ_s^* below which regime change occurs and above which it does not occur is independent of the quality of the agents’ information (recall that this threshold is determined by the comparison between the policy maker’s payoff when, by raising the policy to s , he saves the regime by inducing no agent to attack, and his payoff when, by leaving the policy at $r = \underline{r}$, he induces regime change). In contrast, the threshold $\theta_s^{**}(\sigma)$ above which no type raises the policy to s in any of the equilibria where the range of the policy is $\{\underline{r}, s\}$ is nondecreasing in σ and converges to θ_s^* as $\sigma \rightarrow 0$. The reason is that, as the agents’ information becomes more precise, the size of the attack that the policy maker expects by leaving the policy at $\underline{r}$ converges to zero for all $\theta > \theta_s^*$ (reflecting the fact that all agents correctly realize that regime change will not occur). As a result, in the limit as $\sigma \rightarrow 0$, no type above θ_s^* finds it optimal to raise the policy to s . Along with the fact that there is no equilibrium where some type

below θ_s^* raises the policy to s , independently of the quality of information (recall that raising the policy to s is strictly dominated by leaving the policy to $\underline{r}$ for any $\theta < \theta_s^*$), this means that the measure of types who raise the policy to s vanishes as the quality of the agents' information grows large enough. By implication, a similar result obtains when the econometrician integrates across different equilibria using the selection F .

This last result is interesting because it suggests that the precision of the agents' information need not be important for whether or not regime change occurs in equilibrium³¹, but it may be crucial for whether or not the status quo is maintained with or without policy intervention. Note, however, that this result presumes that the econometrician's beliefs F do not vary with σ . Because the model imposes no relation between F and σ , this is possible, although not necessary.

Remark 5. The result in part (ii) of Proposition 3 delivers a monotonic relation between fundamentals (the policy maker's type) and the regime outcome, while at the same time allowing for some "randomness" in this relation corresponding to the econometrician's uncertainty about the equilibrium being played. This is akin to the predictions one often gets from unique-equilibrium models, except for the following feature. In unique-equilibrium models, the theory typically restricts the residual in the relation between the dependent and the independent variables to be zero; the randomness is then superimposed by the econometrician on the basis of the presumption that there is measurement error or omitted variables that nonetheless do not bias the results. Here, instead, the theory *itself* allows for a random residual: the residual simply captures the econometrician's uncertainty over the equilibrium being played.³²

4.2 Bounds across all selections

As anticipated above, the theory also delivers useful predictions to an econometrician interested in testing, or estimating, the model, but who is not willing to assume (or let the data identify) a particular distribution over the equilibrium being played. This can be done by considering bounds on the probability of regime change and on the probability of intervention *across all possible equilibria* and then investigating how these bounds change with the fundamentals θ (which may be observable ex-post) and/or with the quality of the agents' information σ .

To illustrate, let $D(\theta_1, \theta_2; F, \sigma) \equiv \frac{1}{\theta_2 - \theta_1} \int_{\theta_1}^{\theta_2} D(\theta; F, \sigma) d\theta$ denote the probability of regime change conditional on the event that $\theta \in [\theta_1, \theta_2]$, for given selection $F \in \mathcal{F}(\sigma)$, and $\bar{D}(\theta_1, \theta_2; \sigma) \equiv \sup_{F \in \mathcal{F}(\sigma)} D(\theta_1, \theta_2; F, \sigma)$ and $\underline{D}(\theta_1, \theta_2; \sigma) \equiv \inf_{F \in \mathcal{F}(\sigma)} D(\theta_1, \theta_2; F, \sigma)$ the corresponding bounds across

³¹It may be important only for types $\theta \in (\theta^\#(\sigma), \theta^\#(\sigma'))$ in case one of the pooling equilibria is played, a possibility that is ruled out in the proposition by the assumption that $\text{Supp}[F] \subset (\underline{r}, \lim_{\sigma \rightarrow 0+} r^\#(\sigma))$.

³²Of course, additional randomness can also be superimposed in our model. As anticipated above, all the results in Proposition 3 extend to the case where the regime outcome is affected by random shocks that are exogenous to the interaction between the policy maker and the agents and whose probability is non-increasing in the quality of the fundamentals. The introduction of such shocks makes the relation between the fundamentals and the probability of regime change possibly smoother, without however affecting the qualitative nature of the conclusions.

all selections. Similarly, let $P(r, \theta_1, \theta_2; F, \sigma) \equiv \frac{1}{\theta_2 - \theta_1} \int_{\theta_1}^{\theta_2} P(r, \theta; F, \sigma) d\theta$ denote the probability that the policy is raised to or above r , conditional on the event that $\theta \in [\theta_1, \theta_2]$, for given selection $F \in \mathcal{F}(\sigma)$, and $\bar{P}(r, \theta_1, \theta_2; \sigma) \equiv \sup_{F \in \mathcal{F}(\sigma)} P(r, \theta_1, \theta_2; F, \sigma)$ and $\underline{P}(r, \theta_1, \theta_2; \sigma) \equiv \inf_{F \in \mathcal{F}(\sigma)} P(r, \theta_1, \theta_2; F, \sigma) = 0$ the corresponding bounds. Clearly, $\bar{D}(\theta_1, \theta_2; \sigma) \geq \underline{D}(\theta_1, \theta_2; \sigma)$ and $\bar{P}(r, \theta_1, \theta_2; \sigma) \geq \underline{P}(r, \theta_1, \theta_2; \sigma)$, with strict inequalities when $\underline{\theta} \leq \theta_1 < \theta_2 \leq \theta^\#$ and $\underline{r} < r < r^\#$. That these bounds do not coincide over a subset of the critical region reflects the equilibrium indeterminacy. The next proposition examines how these bounds depend on the fundamentals θ and the quality of information σ .

Proposition 4 (bounds). (i) Fix $\sigma > 0$ and $r \in (\underline{r}, r^\#(\sigma))$. The bounds $\underline{D}(\theta_1, \theta_2; \sigma)$ and $\bar{D}(\theta_1, \theta_2; \sigma)$ are non-increasing in (θ_1, θ_2) , while the bounds $\underline{P}(r, \theta_1, \theta_2; \sigma)$ and $\bar{P}(r, \theta_1, \theta_2; \sigma)$ are non-monotone in (θ_1, θ_2) , in the partial-order sense.

(ii) Fix (θ_1, θ_2) . $\underline{D}(\theta_1, \theta_2; \sigma)$ and $\underline{P}(r, \theta_1, \theta_2; \sigma)$ are independent of σ . For any (θ_1, θ_2) any $\sigma, \sigma' > 0$, $\theta_2 < \min\{\theta^\#(\sigma'), \theta^\#(\sigma)\}$ or $\theta_1 > \max\{\theta^\#(\sigma'), \theta^\#(\sigma)\}$ imply $\bar{D}(\theta_1, \theta_2; \sigma) = \bar{D}(\theta_1, \theta_2; \sigma')$. In contrast, $\lim_{\sigma \rightarrow 0^+} \bar{P}(r, \theta_1, \theta_2; \sigma) = 0$ for any $r > \underline{r}$ and any $\theta_1, \theta_2 \in \mathbb{R}$. In the special case where the agents' gross payoff in case of regime change is fixed (i.e., $Z(\theta, \underline{r}) = z > \underline{r}$ for all θ), then the bound $\bar{D}(\theta_1, \theta_2; \sigma)$ is independent of σ whereas the bound $\bar{P}(r, \theta_1, \theta_2; \sigma)$ is a nondecreasing function of σ .

The results for how the bounds are affected by the fundamentals follow from Proposition 1. Thus consider the effect of information on the bounds. While more precise information need not affect the bounds on the probability of regime change, it affects the bounds on the probability of intervention. In particular, in the limit as $\sigma \rightarrow 0$, the probability of observing policy interventions above r , for any $r > \underline{r}$, vanishes for all measurable sets of θ , whereas the probability of regime change can take any value for any subset of $(\underline{\theta}, \theta^\#(0^+))$, where $\theta^\#(0^+) \equiv \lim_{\sigma \rightarrow 0^+} \theta^\#(\sigma)$. The intuition for these results parallels that for the results in part (iv) of Proposition 3. The key difficulty in establishing these results comes from the need to establish that the equilibrium set converges to its limit-version uniformly across types and selections, as we show in the Appendix.

These properties are particularly sharp when the payoff the agents obtain in case of regime change is independent θ . In this case, both the lower and the upper bound on the probability of regime change are independent of σ ,³³ whereas the upper bound on the probability of policy interventions is decreasing in the quality of information σ^{-1} and vanishes in the limit as $\sigma \rightarrow 0^+$. More generally, what the theory predicts is that policy choices are essentially uniquely determined in the limit, whereas the regime outcomes remain largely indeterminate. We will return to these

³³This follows from the fact that, in this case, the threshold $\theta^\#$ below which regime change occurs and above which it does not occur in any of the pooling equilibria, is independent of σ . It follows that, for any $\theta \in (\underline{\theta}, \theta^\#]$, any $\sigma > 0$, there exist equilibria in which regime change occurs (these are the equilibria that select $\theta_s^* > \theta$) as well as equilibria in which regime change does not occur (these are the equilibria that select $\theta_s^* < \theta$). Hence, the regime outcome remains indeterminate for all $\theta \in (\underline{\theta}, \theta^\#]$. In contrast, as shown in Proposition 3, the probability that the policy is raised to or above any $r > \underline{r}$ vanishes for all positive-measure set of types, a property which is then inherited by the supremum across all selections, using the uniform convergence result mentioned in the main text.

predictions and contrast them to their counterparts under common knowledge in Section 6.

5 Predictions about payoffs

We now turn to the predictions that the theory delivers for the payoff of the policy maker. In contrast to predictions about policy choices and regime outcomes, predictions about payoffs need not be directly testable (the econometrician may not be able to directly observe the policy maker's payoff). Nevertheless, these predictions are important for their policy implications. For example, they permit one to study the ex-ante value that the policy maker may attach to the option to intervene once θ is realized.

For any $s \in [\underline{r}, r^\#]$, let $U_s(\theta; \sigma)$ denote the *lowest* payoff that type θ obtains across all the equilibria in $\mathcal{E}(s; \sigma)$. Next, consider the variant of our model in which r is exogenously fixed at $\underline{r}$ for all θ , interpret this as the game in which the option to intervene is absent, and let $U^o(\theta; \sigma)$ denote the payoff that type θ obtains in the *unique* equilibrium of this game. Clearly, when $s \in \{\underline{r}, r^\#(\sigma)\}$, any equilibrium in $\mathcal{E}(s; \sigma)$ is such that $U_s(\theta; \sigma) = U^o(\theta; \sigma)$ for all θ .³⁴ Thus consider equilibria in which $s \in (\underline{r}, r^\#)$.

Proposition 5 (payoffs). *For any $s \in (\underline{r}, r^\#(\sigma))$, either $U_s(\theta; \sigma) \geq U^o(\theta; \sigma)$ for all θ , with strict inequality for some θ , or there exists a threshold $\theta_s^\dagger(\sigma) \geq \theta^\#(\sigma)$ such that $U_s(\theta; \sigma) < U^o(\theta; \sigma)$ only if $\theta > \theta_s^\dagger(\sigma)$, in which case necessarily $U_s(\theta; \sigma) \geq U^o(\theta; \sigma)$ for all $\theta \leq \theta_s^\dagger(\sigma)$ (with strict inequality if $\theta \in (\theta_s^*, \theta_s^\dagger(\sigma)]$). Moreover, σ small enough ensures that the first case holds.*

To understand this result, note first that types below $\theta^\#$ cannot be worse off with the option to intervene: without this option, they would necessarily experience regime change, whereas with this option they can avoid regime change in some equilibria. Next, consider types above $\theta^\#$. These types can be worse off with the option to intervene only if the size of the attack that they face when they opt for not raising the policy exceeds the one they would have faced absent the option to intervene. In general, this may be possible for some equilibria. However, this possibility vanishes as the precision of information increases. This is because the regime threshold θ_s^* in any of the equilibria with intervention is necessarily lower than $\theta^\#$, the regime threshold in any of the pooling equilibria. Together with the fact that $\theta_s^{**} \rightarrow \theta_s^*$ as $\sigma \rightarrow 0$ guarantees that the size of the attack is also lower, as long as σ is small enough.

The results in Proposition 5 extend to random equilibrium selections. Indeed, fix an arbitrary set of types $[\theta_1, \theta_2] \subset \mathbb{R}$ and an arbitrary selection $F \in \mathcal{F}(\sigma)$ and consider the implied probability that, conditional on $\theta \in [\theta_1, \theta_2]$, the policy maker is strictly worse off with the option to intervene. This probability is zero either for all σ , or at least for σ small enough. Notwithstanding the fact

³⁴The result is immediate when $s = \underline{r}$. For $s = r^\#(\sigma)$, the result follows from the fact that $\theta_s^* = \theta_s^{**}(\sigma) = \theta^\#(\sigma)$ and $X_s(\theta_s^*, \theta_s^{**}(\sigma); \sigma) = x^\#(\sigma)$.

that, in general, the selection F may also depend on σ , this property suggests that the risk of being worse off with the option to intervene vanishes as the agents' information becomes highly precise.

We can also accommodate the case that the selection F changes with σ by considering bounds on equilibrium payoffs across all possible equilibria. Let $\bar{U}(\theta; \sigma)$ and $\underline{U}(\theta; \sigma)$ denote, respectively, the supremum and infimum of the set of equilibrium payoffs that type θ can obtain in the game with the option to intervene, when the quality of information is σ . The following proposition characterizes the relation between these bounds and the payoff obtained in the game in which the option to intervene is absent.

Proposition 6 (payoff bounds). *$\bar{U}(\theta; \sigma) \geq U^o(\theta; \sigma)$ for all $\theta > \underline{\theta}$, with strict inequality for $\theta \in (\underline{\theta}, \theta^\#(\sigma)]$ and with $\lim_{\sigma \rightarrow +\infty} |\bar{U}(\theta; \sigma) - U^o(\theta; \sigma)| = 0$. On the other hand, there exists a threshold $\theta^\dagger(\sigma) \geq \theta^\#(\sigma)$ such that $\underline{U}(\theta; \sigma) < U^o(\theta; \sigma)$ only if $\theta > \theta^\dagger(\sigma)$. Finally, $\lim_{\sigma \rightarrow 0+} \underline{U}(\theta; \sigma) = \lim_{\sigma \rightarrow 0+} U^o(\theta; \sigma)$ for all θ .*

The result that $\bar{U}(\theta; \sigma) \geq U^o(\theta; \sigma)$ follows from the fact that, in the game with the option to intervene, there always exist equilibria where each type above $\theta > \underline{\theta}$ can induce all agents to not attack by raising the policy infinitesimally above $\underline{r}$ thus obtaining a payoff arbitrarily close to the highest *feasible* payoff $W(\theta, \underline{r}, 0)$. In contrast, in the game without the option to intervene, the highest feasible payoff $W(\theta, \underline{r}, 0)$ may be attainable only by sufficiently high types (for whom $A(\theta, \underline{r})$ becomes arbitrary small). These observations explain the results in the proposition that pertain to the upper bound on the equilibrium payoff. The results for the lower bounds follow from essentially the same arguments discussed above in relation to Proposition 5—very high types can be worse off with the option to intervene only if there exist equilibria in which the size attack that each type faces when he does not raise the policy above $\underline{r}$ is larger than in the game without the option to intervene. However, even if such equilibria exist, low types continue to be (weakly) better off with the option to intervene, no matter the equilibrium being played, for these types would have experienced regime change without the option to intervene. Finally, the result that, as $\sigma \rightarrow 0$, $\underline{U}(\theta; \sigma)$ converges to $U^o(\theta; \sigma)$ for all θ follows from the fact that, as $\sigma \rightarrow 0$, the lower bound on the equilibrium payoffs in the game with the option to intervene is attained under any of the pooling equilibria, which is clearly the same payoff as in the game in which the option to intervene is absent.

Now imagine that, before knowing his type, the policy maker decides whether to maintain or to give up the option to intervene after learning θ . The aforementioned results suggest that, in general, the policy maker need not be able to ensure that he will be better off with the option to intervene no matter the realized θ : he may get “trapped” in an equilibrium in which he is worse off when θ turns out to be sufficiently high. Even then, however, the policy maker is better off for low θ . Therefore, the option to intervene either is beneficial for all θ , or it implements a form of ex-ante insurance across types.

6 Contrast to common knowledge

We now contrast the predictions that the theory delivers for the incomplete-information game with those for its common-knowledge counterpart. We further show that, while multiplicity obtains in our model for any level of noise, the set of equilibrium outcomes becomes smaller and smaller (in an appropriate sense) as the quality of information improves—but it explodes when the noise is zero. The purpose of these exercises is two-fold: (i) to highlight that the selection power of global games retains significant bite also in games like ours where the endogeneity of information leads to multiple equilibria; and (ii) to establish that the predictions that we have identified, albeit intuitive, would not have been possible with complete information.

Recalling that $\rho(\theta)$ denotes the maximal level of r that is not strictly dominated by $\underline{r}$ for a policy maker of type θ , we have the following result.

Proposition 7 (common knowledge). *Consider the game with $\sigma = 0$.*

(i) *A policy $r(\cdot)$ can be part of a subgame-perfect equilibrium if and only if $r(\theta) \leq \rho(\theta)$ for $\theta \in (\underline{\theta}, \bar{\theta}]$ and $r(\theta) = \underline{r}$ for $\theta \notin (\underline{\theta}, \bar{\theta}]$.*

(ii) *A regime outcome $D(\cdot)$ can be part of a subgame-perfect equilibrium if and only if $D(\theta) = 1$ for $\theta \leq \underline{\theta}$, $D(\theta) \in \{0, 1\}$ for $\theta \in (\underline{\theta}, \bar{\theta}]$, and $D(\theta) = 0$ for $\theta > \bar{\theta}$.*

This result contrasts sharply with the results in Propositions 1-4. None of the predictions in the game with incomplete information are valid in the game with common knowledge. In particular, the policy can now take any shape in the critical region $(\underline{\theta}, \bar{\theta}]$. Similarly, the probability of regime change can take any value within the critical region and need not be monotone in θ . In essence, “almost anything goes” within the critical region under complete information. In particular, the only policy choices and regime outcomes that are ruled out by equilibrium reasoning under complete-information are those that are ruled out by strict dominance. A similar “anything-goes” result holds if one looks at the policy maker’s payoff.

The contrast between the complete- and incomplete-information versions of our model is most evident in the limit as $\sigma \rightarrow 0^+$. Let $\mathcal{G}(\sigma)$ denote the set of pairs (θ, r) such that, in the game with noise $\sigma \geq 0$, there is an equilibrium in which type θ sets the policy at r . Next, let $\underline{\rho}^+ \equiv \lim_{\theta \rightarrow \underline{\theta}^+} \rho(\theta)$.³⁵ We then have the following result.

Proposition 8 (limit outcomes). *Under complete information,*

$$\mathcal{G}(0) = \{(\theta, r) : \text{either } \theta \in (\underline{\theta}, \bar{\theta}] \text{ and } \underline{r} \leq r \leq \rho(\theta), \text{ or } \theta \notin (\underline{\theta}, \bar{\theta}] \text{ and } r = \underline{r}\}.$$

In contrast, under incomplete information,

$$\lim_{\sigma \rightarrow 0^+} \mathcal{G}(\sigma) = \left\{(\theta, r) : \text{either } r = \underline{r} \text{ and } \theta \in \mathbb{R}, \text{ or } r \in (\underline{\rho}^+, r^\#(0^+)] \text{ and } \theta = \rho^{-1}(\theta)\right\},$$

which is a zero-measure subset of $\mathcal{G}(0)$.

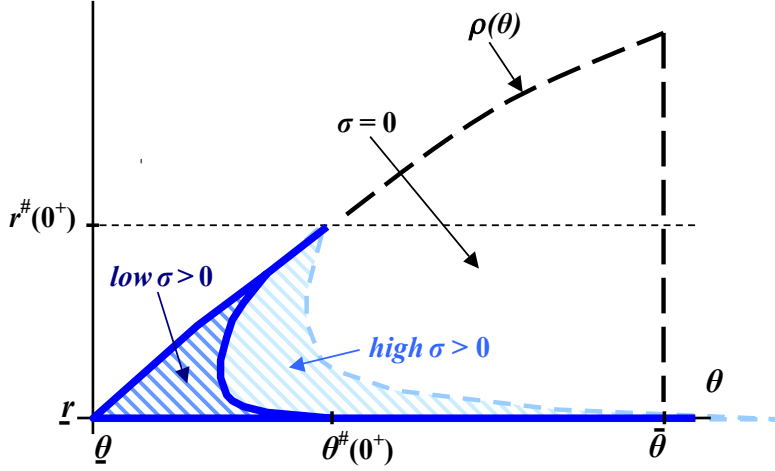


Figure 1: Pairs (θ, r) that can be observed as equilibrium outcomes.

This result, which follows from the fact that, as $\sigma \rightarrow 0^+$, $\theta_s^{**} \rightarrow \theta_s^*$ for all $s > \underline{r}$, is illustrated in Figure 1 for an example in which $Z(\theta, r) = z$, $W(\theta, r, A) = V(\theta, A) - r$ and $L(\theta, r) = -r$ and $R(\theta, r, A) = V(\theta, A)$. The common-knowledge set, $\mathcal{G}(0)$, is given by the large triangular area. The incomplete-information set, $\mathcal{G}(\sigma)$ for $\sigma > 0$ is given by the dashed blue area. In this case, as long as $\sigma > 0$, the lower σ , the smaller the set of policies that can be played by any given θ , and hence the smaller the dashed area in Figure 1 (i.e., $\sigma' > \sigma > 0$ implies $\mathcal{G}(\sigma') \subset \mathcal{G}(\sigma)$). The monotonicity of $\mathcal{G}(\sigma)$ in σ does not necessarily hold in the case where the agents' payoff in case of regime change depends θ . However, as the proposition makes clear, what is true more generally is that, in the limit, as the noise in information vanishes, the set of policies that can be sustained in equilibrium for almost any given θ is a zero measure subset of the set of policies that can be sustained under common knowledge. More precisely, the set $\mathcal{G}(\sigma)$ converges to the boundary points of the set of policies that would have been possible under complete information for almost any $\theta \leq \theta^\#(0^+) \equiv \lim_{\sigma \rightarrow 0^+} \theta^\#(\sigma)$, and to the cost-minimizing policy $\underline{r}$ for $\theta > \theta^\#(0^+)$.

7 Concluding Remarks

The approach followed in most applications of global games is to use incomplete information as a tool to select a unique equilibrium in coordination settings that admit multiple equilibria under common knowledge—to assume certain exogenous information structures that ensure uniqueness, without investigating what determines information in the first place. For certain questions, however, the endogeneity of information is central to the phenomenon under examination and often brings

³⁵Note that, in general, $\underline{\rho}^+$ can be strictly higher than $\rho(\underline{\theta}) = \underline{r}$.

back multiple equilibria. The broader methodological contribution of this paper is to illustrate that this multiplicity may be very different from the one that obtains under common knowledge and, more importantly, that it need not preclude concrete, intuitive, and testable predictions.

These predictions typically take the form of stochastic non-linear relations between the endogenous variables of interest (e.g., the probability of regime change and of policy interventions) and the primitives of the model (e.g., the underlying fundamentals) with the randomness originating in the econometrician's uncertainty over the equilibrium being played as opposed to measurement error (as in unique-equilibrium models). Clearly, confronting these predictions with the data remains more challenging than in unique-equilibrium models. Yet, recent advances in structural estimation of multiple-equilibria models (e.g., Tamer, 2003, Ciliberto and Tamer, 2009, Grieco, 2010) may help in this direction.

The procedure of iterated deletion of non-equilibrium strategies we used in the present paper to identify sharp predictions rests, of course, on the specific channel of information endogeneity we considered (signaling). In future work, it would be interesting to investigate how analogous procedures can be used to deliver tight predictions in games where the endogeneity of information comes from alternative channels, such as the aggregation of information through prices, or social learning in dynamic settings.³⁶

³⁶Complementary in this respect is the recent paper by Chassang (2010). That paper applies global-game techniques in a dynamic exit game in which repeated play sustains multiple equilibria despite the incompleteness of information. The origin of multiplicity in that paper is therefore very different from the one that is the focus of our paper. The two papers are nevertheless complementary in the sense that they both show how global-game techniques can retain significant selection power even when they fail to deliver a unique equilibrium.

Appendix A: Additional results

Proposition 1 in the main text left open the possibility of equilibria in which the policy is raised to some $s > r^\#$. This possibility was ruled out in Proposition 2 by imposing *SCC*. Part (i) in the next proposition shows that this possibility vanishes in the limit as $\sigma \rightarrow 0$ even when *SCC* does not hold, while parts (ii) and (iii) study the limit behavior of the range of policy intervention.

Proposition 9 (necessary conditions for small σ). *For any $\varepsilon > 0$, there exists $\bar{\sigma} > 0$ such that the following results hold for any $\sigma < \bar{\sigma}$:*

- (i) $\mathcal{E}(s; \sigma) \neq \emptyset$ only if $s \leq r^\#(\sigma) + \varepsilon$;
- (ii) whenever $\mathcal{E}(s; \sigma) \neq \emptyset$ then $\theta_s^* \leq \theta^\#(\sigma) + \varepsilon$;
- (iii) whenever $\mathcal{E}(s; \sigma) \neq \emptyset$ and $s \geq \underline{r} + \varepsilon$ then $\theta_s^{**}(\sigma) \leq \theta_s^* + \varepsilon$.

Proposition 2 also left open the possibility of equilibria with $s > \underline{r}$ in which the policy is raised for only a strict subset of the interval $(\theta_s^*, \theta_s^{**})$. With *SCC*, this possibility rests on the agents following a non-monotone strategy. This possibility can be ruled out by letting the distribution of the noise be log-concave—a restriction that helps guaranteeing that, along any equilibrium, the agents’ posteriors about θ , and hence their strategies, are monotone in their signals.

Proposition 10 (monotonicity of speculators’ behavior). (i) *Suppose the noise distribution ψ is log-concave. Then, for any $\sigma > 0$, any $s \in (\underline{r}, r^\#(\sigma)]$, $\mathcal{E}(s; \sigma) \neq \emptyset$.*

(ii) *Suppose *SCC* holds and ψ is log-concave. Then for any $\sigma > 0$ and any $s \in (\underline{r}, r^\#(\sigma)]$, any equilibrium in $\mathcal{E}(s; \sigma)$ is such that $r(\theta) = s$ for all $\theta \in (\theta_s^*, \theta_s^{**})$.*

Part (i) in the above proposition establishes that log-concavity of ψ is, by itself, another sufficient condition for the existence of equilibria in which the policy maker raises the policy to s , for any $s \in (\underline{r}, r^\#(\sigma)]$.³⁷ Part (ii) then establishes that the combination of log-concavity of ψ with *SSC* suffices for ruling out equilibria in which the policy is raised for only a subset of the interval $(\theta_s^*, \theta_s^{**})$, which further sharpens the interpretation of $(\theta_s^*, \theta_s^{**})$.

Appendix B: Omitted proofs

Proof of Lemma 1. Considers the continuation game that follows the observation of $\underline{r}$. Because the observation of $\underline{r}$ conveys no information, this game is essentially a standard global game (e.g., Morris and Shin, 1998, 2003). The proof below shows that this game admits a unique continuation

³⁷The role of log-concavity is to guarantee that, irrespective of the shape of the policy r in the region $[\theta_s^*, \theta_s^{**}]$, the probability that each agent assigns to regime change when observing no intervention is necessarily decreasing in the signal x . This in turn implies that the size of the attack $A(\theta, \underline{r})$ that the policy maker expects when he does not intervene is necessarily decreasing in θ . As we show in the proof in Appendix B, under such monotonicities, one can construct equilibria in which intervention occurs for a (possibly non-connected) subset of $[\theta_s^*, \theta_s^{**}]$.

equilibrium in monotone strategies. Standard results from global games (based on iterated deletion of strictly dominated strategies) then imply that this equilibrium is the unique equilibrium of the continuation game.

Clearly any monotone continuation equilibrium must be characterized by thresholds $x^\#(\sigma)$ and $\theta^\#(\sigma)$ such that all agents attack if $x < x^\#(\sigma)$ and not attack if $x > x^\#(\sigma)$ in which case regime change occurs if $\theta \leq \theta^\#(\sigma)$ and does not occur if $\theta > \theta^\#(\sigma)$. Hereafter, we show that such a continuation equilibrium exists and is unique. To simplify the notation, we momentarily drop the dependence of the thresholds $x^\#(\sigma)$ and $\theta^\#(\sigma)$ on σ .

First, note that an agent with signal x who expects regime change to occur if and only if $\theta \leq \theta^\#$ finds it optimal to attack if and only if

$$\int_{-\infty}^{\theta^\#} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{x-\tilde{\theta}}{\sigma}\right) d\tilde{\theta} - Q(\underline{r}) \geq 0. \quad (13)$$

Because $Z(\cdot, \underline{r})$ is non-increasing, the left-hand-side of (13) is continuously strictly decreasing in x .³⁸ Furthermore, it is strictly positive for x small enough and strictly negative for x large enough. It follows that the inequality in (13) holds if and only if $x \leq x^\#$ where $x^\#$ is the unique solution to

$$\int_{-\infty}^{\theta^\#} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{x^\#-\tilde{\theta}}{\sigma}\right) d\tilde{\theta} = Q(\underline{r}). \quad (14)$$

Next, consider the fate of the regime. Because $A(\theta, \underline{r}) = \Psi((x^\# - \theta)/\sigma)$ is decreasing in θ , regime change occurs if and only if $\theta \leq \theta^\#$ where the threshold $\theta^\#$ is the unique solution to

$$R\left(\theta^\#, \underline{r}, \Psi\left(\frac{x^\#-\theta^\#}{\sigma}\right)\right) = 0. \quad (15)$$

Thus any monotone equilibrium is identified by a solution $(x^\#, \theta^\#)$ to conditions (14) and (15). Below, we show that a solution to these conditions exists and is unique.

To this aim, let $\theta^\#(x^\#)$ be the implicit function defined by (15). By the Implicit Function Theorem,

$$\frac{d\theta^\#(x^\#)}{dx^\#} = \frac{-R_A\left(\theta^\#, \underline{r}, \Psi\left(\frac{x^\#-\theta^\#}{\sigma}\right)\right) \psi\left(\frac{x^\#-\theta^\#}{\sigma}\right)}{-R_A\left(\theta^\#, \underline{r}, \Psi\left(\frac{x^\#-\theta^\#}{\sigma}\right)\right) \psi\left(\frac{x^\#-\theta^\#}{\sigma}\right) + \sigma R_\theta\left(\theta^\#, \underline{r}, \Psi\left(\frac{x^\#-\theta^\#}{\sigma}\right)\right)} \in (0, 1).$$

Next, let $LHS(x^\#)$ denote the function of $x^\#$ that is defined by the left-hand side of (14) once we replace $\theta^\#$ with $\theta^\#(x^\#)$. Differentiating with respect to $x^\#$ gives the following expression:

$$\begin{aligned} \frac{\partial LHS(x^\#)}{\partial x^\#} &= \frac{d\theta^\#(x^\#)}{dx^\#} Z\left(\theta^\#(x^\#), \underline{r}\right) \frac{1}{\sigma} \psi\left(\frac{x^\#-\theta^\#(x^\#)}{\sigma}\right) + \\ &\quad + \int_{-\infty}^{\theta^\#(x^\#)} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma^2} \psi'\left(\frac{x^\#-\tilde{\theta}}{\sigma}\right) d\tilde{\theta}. \end{aligned} \quad (16)$$

³⁸Note that, in any pooling equilibrium, after observing $r = \underline{r}$, the agents' posteriors beliefs that $\tilde{\theta} < \theta$ are given by $1 - \Psi((x - \theta)/\sigma)$. These beliefs can be ranked according to FOSD, implying that, given any non-increasing function $f(\theta)$, the expected value of f given x is decreasing in x .

Integrating by parts, the last term in (16) can be rewritten as

$$\begin{aligned} & -Z\left(\theta^\#(x^\#), \underline{r}\right) \frac{1}{\sigma} \psi\left(\frac{x^\# - \theta^\#(x^\#)}{\sigma}\right) + \lim_{\tilde{\theta} \rightarrow -\infty} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{x^\# - \tilde{\theta}}{\sigma}\right) \\ & + \int_{-\infty}^{\theta^\#(x^\#)} Z_\theta\left(\tilde{\theta}, \underline{r}\right) \frac{1}{\sigma} \psi\left(\frac{x^\# - \tilde{\theta}}{\sigma}\right) d\tilde{\theta}. \end{aligned}$$

It follows that (16) reduces to

$$\begin{aligned} \frac{\partial LHS(x^\#)}{\partial x^\#} &= \left(\frac{d\theta^\#(x^\#)}{dx^\#} - 1\right) Z\left(\theta^\#(x^\#), \underline{r}\right) \frac{1}{\sigma} \psi\left(\frac{x^\# - \theta^\#(x^\#)}{\sigma}\right) \\ &+ \int_{-\infty}^{\theta^\#(x^\#)} Z_\theta\left(\tilde{\theta}, \underline{r}\right) \frac{1}{\sigma} \psi\left(\frac{x^\# - \tilde{\theta}}{\sigma}\right) d\tilde{\theta} \end{aligned}$$

which is clearly negative. Lastly, note that

$$\lim_{x^\# \rightarrow -\infty} LHS(x^\#) > Q(\underline{r}) > 0 = \lim_{x^\# \rightarrow +\infty} LHS(x^\#)$$

By the Intermediate Value Theorem, it then follows that there exists a unique $x^\#$ such that $LHS(x^\#) = Q(\underline{r})$. Given the uniqueness of $x^\#$, the uniqueness of $\theta^\#$ follows immediately from (15). This establishes existence and uniqueness of a monotone continuation equilibrium. *Q.E.D.*

Proof of Lemma 2. Because the policy maker faces no uncertainty about the size of the attack, he can predict the fate of the regime with certainty. The assumption that the payoff L in case of regime change is independent of A then implies that raising the policy above $\underline{r}$ and experiencing regime change yields a lower payoff than leaving the policy at $\underline{r}$.³⁹ It follows that any type who in equilibrium raises the policy above $\underline{r}$ must be spared from regime change, for otherwise he would be strictly better off by setting $r = \underline{r}$. By implication, the observation of any equilibrium policy $r > \underline{r}$ necessarily signals to the agents that the status quo will be maintained and thus induces each agent to not attack no matter his signal x . But then any type of the policy maker can always save on the cost of intervention by setting the lowest $r > \underline{r}$ among those that are played in equilibrium. Hence in any equilibrium in which some type intervenes, there exists a single $s > \underline{r}$ such that $r(\theta) = s$ whenever $r(\theta) \neq \underline{r}$. Furthermore, no speculator attacks following the observation of $r = s$ which means that $A(\theta, s) = 0$ for all θ . *Q.E.D.*

Proof of Lemma 3. That $\mathcal{E}(s; \sigma) = \emptyset$ for any $s > \underline{r}$ for which $W(\theta, s, 0) < L(\theta, \underline{r})$ for all θ follows directly from the fact that, in this case, the net payoff that each type θ obtains by raising the policy to s is strictly less than the payoff that the same type obtains by leaving the policy at $\underline{r}$

³⁹When leaving the policy at $\underline{r}$ leads to regime change, it follows from the assumption that $L(\theta, r)$ is strictly decreasing in r ; when, instead, it leads to no regime change, it follows from the assumption that $W(\theta, \underline{r}, A) > L(\theta, \underline{r})$ when $R(\theta, \underline{r}, A) > 0$.

(this is, irrespective of the agents' behavior and of whether or not leaving the policy at $\underline{r}$ leads to regime change). Thus consider s for which there exists a $\theta \geq \underline{\theta}$ such that

$$W(\theta, s, 0) \geq L(\theta, \underline{r}).$$

The assumption that $W(\theta, s, 0) - L(\theta, \underline{r})$ is strictly increasing in θ then implies that any type $\theta > \theta_s^*$, by setting $r = s$, can guarantee himself a payoff strictly higher than the payoff that the same type obtains by setting $r = \underline{r}$ and experiencing regime change (recall that, from Lemma 2, $A(\theta, s) = 0$ for any θ). But then, in equilibrium, no type above θ_s^* experiences regime change. *Q.E.D.*

Proof of Lemma 4. Consider the sequence $\{\theta^n\}_{n=0}^\infty$ constructed as follows. First, let

$$\theta^0 \equiv \inf \{ \theta \geq \theta_s^* : W(\theta, s, 0) < W(\theta, \underline{r}, 1) \text{ and } R(\theta, \underline{r}, 1) > 0 \}.$$

That $\theta^0 < +\infty$ follows from the assumption that $\lim_{\theta \rightarrow +\infty} \{W(\theta, s, 0) - W(\theta, \underline{r}, 1)\} < 0$, which simply imposes that raising the policy is dominated for sufficiently high types. Next, for any $n \geq 1$, let

$$\theta^n \equiv \inf \{ \theta \geq \theta_s^* : g(\theta; \theta_s^*, \theta^{n-1}, \sigma) < 0 \}$$

where, for any $\theta \geq \underline{\theta}$, any (θ_s^*, θ') with $\theta' \geq \theta_s^*$, and any $\sigma > 0$, the function g is defined by

$$g(\theta; \theta_s^*, \theta', \sigma) \equiv U(\theta, s, 0) - U(\theta, \underline{r}, \bar{A}(\theta; \theta_s^*, \theta'))$$

where $\bar{A}(\theta; \theta_s^*, \theta') \equiv \Psi\left(\frac{X(\theta_s^*, \theta'; \sigma) - \theta}{\sigma}\right)$ and where $X(\theta_s^*, \theta'; \sigma)$ is the unique solution to (5). The function $g(\theta; \theta_s^*, \theta', \sigma)$ thus identifies the differential between the payoff that type θ obtains by raising the policy to $r = s$, facing no attack and avoiding regime change, and the payoff that the same type obtains by leaving the policy at $r = \underline{r}$, facing an attack of size $\bar{A}(\theta; \theta_s^*, \theta')$ —that is, the attack implied by the agents play according to the threshold strategy with cutoff $X(\theta_s^*, \theta'; \sigma)$ —and then facing regime change if and only if $R(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta'; \sigma) - \theta}{\sigma}\right)) \leq 0$. Note that, because $\theta_s^* \geq \underline{\theta}$, $U(\theta, s, 0) = W(\theta, s, 0)$ for all $\theta > \theta_s^*$. Furthermore, because $W(\theta, s, 0) > L(\theta, \underline{r})$ for all $\theta > \theta_s^*$, then $g(\theta; \theta_s^*, \theta', \sigma) < 0$ only if $R(\theta, \underline{r}, \bar{A}(\theta; \theta_s^*, \theta')) > 0$; that is, for any type above θ_s^* to be better off by leaving the policy at $r = \underline{r}$, it must be that the size of the attack $\bar{A}(\theta; \theta_s^*, \theta')$ triggered by leaving the policy at $\underline{r}$ does not lead to regime change.

This sequence has a simple interpretation. In any equilibrium in which the range of the policy is $r(\mathbb{R}) = \{\underline{r}, s\}$, no type $\theta \notin [\theta_s^*, \theta^0]$ raises the policy to $r = s$. Given so, an agent who expects regime change to occur if and only if $\theta < \theta_s^*$ and $r(\theta) = \underline{r}$ if and only if $\theta \notin [\theta_s^*, \theta^0]$ finds it optimal to attack, when observing $\underline{r}$, if and only if he receives a signal $x \leq X(\theta_s^*, \theta^0; \sigma)$. To see this, recall that $X(\theta_s^*, \theta^0; \sigma)$ is the unique solution to

$$\frac{\int_{-\infty}^{\theta_s^*} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{x - \tilde{\theta}}{\sigma}\right) d\tilde{\theta}}{1 - \Psi\left(\frac{x - \theta_s^*}{\sigma}\right) + \Psi\left(\frac{x - \theta^0}{\sigma}\right)} - Q(\underline{r}) = 0 \quad (17)$$

Along with the fact that the expected payoff from attacking (the left hand side of (17)) is positive for sufficiently low x and negative for sufficiently high x , this gives the result.

But then, by implication, an agent who expects the status quo to be maintained for all $\theta \geq \theta_s^*$ (but possibly also for some $\theta < \theta_s^*$) and $r(\theta) = \underline{r}$ for all $\theta \notin [\theta_s^*, \theta^0]$ (but possibly also for some $\theta \in [\theta_s^*, \theta^0]$) never finds it optimal to attack for $x > X(\theta_s^*, \theta^0; \sigma)$. To see this, note that when the status quo is maintained also for some $\theta < \theta_s^*$, the payoff that the agent expects from attacking when he observes $\underline{r}$ is smaller than when regime change occurs for *all* $\theta < \theta_s^*$. Similarly, when the policy maker leaves the policy at $\underline{r}$ also for some $\theta \in [\theta_s^*, \theta^0]$, the observation $r = \underline{r}$ maps to a lower posterior probability of regime change than when $r = s$ for *all* $\theta \in [\theta_s^*, \theta^0]$. Hence, the incentives to attack after observing $\underline{r}$ are maximal when regime change occurs for all $\theta < \theta_s^*$ and $r(\theta) = s$ for all $\theta \in [\theta_s^*, \theta^0]$, which explains why no agent ever finds it optimal to attack for $x > X(\theta_s^*, \theta^0; \sigma)$. Knowing this, a policy maker who expects no agent to attack for $x > X(\theta_s^*, \theta^0; \sigma)$ never finds it optimal to raise the policy to $r = s$ for any $\theta > \theta^1$. Knowing this, no agent finds it optimal to attack for any $x > X(\theta_s^*, \theta^1; \sigma)$ when observing $\underline{r}$, and so on.

Below, we conclude the proof by establishing that the sequence $\{\theta^n\}_{n=0}^\infty$ is non-increasing. Because it is also bounded from below by θ_s^* , it has to converge. Now the limit is $\theta_s^{**} = \max\{\hat{\theta}_s, \dot{\theta}_s\}$ where

$$\dot{\theta}_s = \sup\{\theta \geq \theta_s^* : W(\theta, s, 0) = W(\theta, \underline{r}, B(\theta_s^*, \theta; \sigma)) \text{ and } R(\theta, \underline{r}, B(\theta_s^*, \theta; \sigma)) > 0\}.$$

$$\hat{\theta}_s \equiv \sup\{\theta \geq \theta_s^* : R(\theta, \underline{r}, B(\theta_s^*, \theta; \sigma)) = 0\}$$

if at least one of the above two sets is non-empty. Else, the limit is θ_s^* . In the latter case, by the definition of the sequence $\{\theta^n\}_{n=0}^\infty$, no type above θ_s^* is willing to raise the policy to $r = s$. Together with the fact that $W(\theta, s, 0) < L(\theta, \underline{r})$ for any $\theta < \theta_s^*$, meaning that no type below θ_s^* is also willing to raise the policy, then this means that $\mathcal{E}(s; \sigma) = \emptyset$. We conclude that, if $\mathcal{E}(s; \sigma) \neq \emptyset$, there must exist a $\theta_s'' \geq \theta_s^*$ such that (i) either $W(\theta_s'', s, 0) = W(\theta_s'', \underline{r}, B(\theta_s^*, \theta_s''; \sigma))$ and $R(\theta_s'', \underline{r}, B(\theta_s^*, \theta_s''; \sigma)) > 0$, or (ii) $R(\theta_s'', \underline{r}, B(\theta_s^*, \theta_s''; \sigma)) = 0$.

We conclude the proof by showing that the sequence $\{\theta^n\}_{n=0}^\infty$ defined above is indeed non-increasing. To see this, fix any $s > \underline{r}$ for which there exists a $\theta \geq \underline{\theta}$ such that $W(\theta, s, 0) \geq L(\theta, \underline{r})$ and let $\theta_s^* = \inf\{\theta \geq \underline{\theta} : W(\theta, s, 0) \geq L(\theta, \underline{r})\}$. Towards a contradiction, suppose that there exists an $n \geq 1$ such that $\theta^n > \theta^{n-1}$. Without loss of generality, then let $n \geq 1$ be the first step in the sequence for which $\theta^n > \theta^{n-1}$ (i.e., $\theta^j \leq \theta^{j-1}$ for all $j \leq n-1$). By the definition of θ^{n-1} , then for any $\theta > \theta^{n-1}$,

$$g(\theta; \theta_s^*, \theta^{n-2}, \sigma) < 0.$$

Because $W(\theta, s, 0) > L(\theta, \underline{r})$ for any $\theta > \theta_s^*$, this means that, for any $\theta > \theta^{n-1}$, necessarily $R(\theta, \underline{r}, \Psi(\frac{X(\theta_s^*, \theta^{n-2}) - \theta}{\sigma})) > 0$, meaning that the policy maker can avoid regime change by leaving

the policy at $\underline{r}$. This in turn implies that, for any $\theta > \theta^{n-1}$, the payoff by not raising the policy is

$$W\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta^{n-2}) - \theta}{\sigma}\right)\right)$$

Because $\theta^{n-1} \leq \theta^{n-2}$ and because $X(\theta_s^*, \cdot; \sigma)$ is increasing, this in turn implies that, for any $\theta > \theta^{n-1}$,

$$g(\theta; \theta_s^*, \theta^{n-1}, \sigma) < 0.$$

By the definition of θ^n , this means that $\theta^n \leq \theta^{n-1}$, which proves that the sequence $\{\theta^n\}_{n=0}^\infty$ is non-increasing. *Q.E.D.*

Proof of Lemma 5. That $D(\theta) = 1$ for any $\theta < \theta_s^*$ is trivial when $\theta_s^* = \underline{\theta}$. Thus assume $\theta_s^* > \underline{\theta}$. The result that $D(\theta) = 1$ for any $\theta < \min\{\theta_s^*, \theta^\#(\sigma)\}$ is then established by comparing the agents' incentives to attack after observing $\underline{r}$ with the corresponding incentives when they expect $r(\theta) = \underline{r}$ for all θ .

Let $\{\theta_n, x_n\}_{n=0}^\infty$ be the following sequence. First, let $\theta_0 \equiv \underline{\theta}$ and let x_0 be implicitly defined by

$$\int_{-\infty}^{\underline{\theta}} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{x_0 - \tilde{\theta}}{\sigma}\right) d\tilde{\theta} = Q(\underline{r}).$$

Next, for any $n \geq 1$, let $\theta_n \equiv \min\{\theta_s^*, \theta'_n\}$, where θ'_n solves $R(\theta'_n, \underline{r}, \Psi(\frac{x_{n-1} - \theta'_n}{\sigma})) = 0$, and let x_n be implicitly defined by

$$\int_{-\infty}^{\theta'_n} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{x_n - \tilde{\theta}}{\sigma}\right) d\tilde{\theta} = Q(\underline{r}).$$

This sequence also has a simple interpretation. An agent who observes $r = \underline{r}$ and believes that $r(\theta) = \underline{r}$ for all θ and that no other agent attacks (in which case regime change occurs if and only if $\theta \leq \underline{\theta}$), finds it optimal to attack if and only if $x \leq x_0$. By implication, an agent who expects no other agent to attack and $r(\theta) = \underline{r}$ for all $\theta < \theta_s^*$ (but possibly $r(\theta) > \underline{r}$ for some $\theta > \theta_s^*$) necessarily finds it optimal to attack for any $x < x_0$. This simply follows from the fact that the observation of $\underline{r}$ is most informative of regime change when all types for whom regime change occurs set $r = \underline{r}$, while some of the types who save the regime raise the policy above $\underline{r}$. However, if all agents attack whenever $x < x_0$, regime change occurs for all $\theta < \theta_1$. This in turn implies that there exists an $x_1 > x_0$ such that an agent who expects all other agents to attack if $x < x_0$ (and hence regime change to occur for all $\theta \leq \theta_1$) and who believes that $r(\theta) = \underline{r}$ for all θ , necessarily finds it optimal to attack for all $x < x_1$. By implication, an agent who expects $r(\theta) = \underline{r}$ for all $\theta < \theta_s^*$ but possibly $r(\theta) > \underline{r}$ for some $\theta > \theta_s^*$, necessarily finds it optimal to attack for any $x < x_1$, and so on.

Because $\{\theta_n\}_{n=0}^\infty$ is increasing and bounded from above it necessarily converges. Note that R and Ψ are continuous and that the unique solution (x, θ') to $R(\theta', \underline{r}, \Psi(\frac{x - \theta'}{\sigma})) = 0$ and

$$\int_{-\infty}^{\theta'} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{x - \tilde{\theta}}{\sigma}\right) d\tilde{\theta} = Q(\underline{r})$$

is attained at $\theta' = \theta^\#(\sigma)$ and $x = x^\#(\sigma)$ (by uniqueness of the monotone equilibrium established in Lemma 1). It follows that $\lim_{n \rightarrow +\infty} \theta_n = \theta_s^*$ if $\theta_s^* \leq \theta^\#(\sigma)$ and $\lim_{n \rightarrow +\infty} \theta_n = \theta^\#(\sigma)$ otherwise. By implication, $D(\theta) = 1$ for all $\theta < \min\{\theta_s^*, \theta^\#(\sigma)\}$. This establishes the first part of the Lemma.

That $\theta_s^* > \theta^\#(\sigma)$ if and only if $s > r^\#(\sigma)$ follows from the fact that, for any θ , $W(\theta, \cdot, 0) - L(\theta, \underline{r})$ is continuous and strictly decreasing in s , while for any $s > \underline{r}$, $W(\cdot, s, 0) - L(\cdot, \underline{r})$ is continuous and strictly increasing in θ . This establishes the second part of the lemma. *Q.E.D.*

Proof of Proposition 2. We start by establishing existence of pooling equilibria.

Lemma A0. *For any $\sigma > 0$, $\mathcal{E}(\underline{r}; \sigma) \neq \emptyset$.*

Proof of Lemma A0. Fix $\sigma > 0$. We establish the result by proving existence of a strategy profile for the agents, along with a system of supporting beliefs, such that, given this profile, no type of the policy maker finds it optimal to raise the policy above $\underline{r}$. Consider the following strategy profile for the agents: (i) for $r = \underline{r}$, $a(x, r) = 1$ if and only if $x < x^\#(\sigma)$ where $x^\#(\sigma)$ is the unique threshold defined in the proof of Lemma 1; for any $r \in (\underline{r}, \rho(\bar{\theta})]$, $a(x, r) = 1$ irrespective of x ; for any $r > \rho(\bar{\theta})$, $a(x, r) = 0$ irrespective of x .

Given this profile, no type of the policy maker finds it optimal to raise the policy. This is immediate for any $\theta \leq \bar{\theta}$: for these types, raising the policy to $r \in (\underline{r}, \rho(\bar{\theta})]$ leads to regime change (with payoff $L(\theta, r) < L(\theta, \underline{r})$) whereas raising the policy above $\rho(\bar{\theta})$ is strictly dominated by leaving the policy to $\underline{r}$. Thus consider a type $\theta > \bar{\theta}$. Clearly, raising the policy to $r \in (\underline{r}, \rho(\bar{\theta})]$ yields a lower payoff than leaving the policy at $\underline{r}$, for it is costly and it increases the measure of agents attacking. That raising the policy to $r > \rho(\bar{\theta})$ also yields a lower payoff than leaving the policy at $\underline{r}$ follows from the fact that

$$W(\theta, r, 0) < W(\theta, \rho(\bar{\theta}), 0) < W\left(\theta, \underline{r}, \Psi\left(\frac{x^\#(\sigma) - \bar{\theta}}{\sigma}\right)\right) \leq W\left(\theta, \underline{r}, \Psi\left(\frac{x^\#(\sigma) - \theta}{\sigma}\right)\right)$$

where the first inequality follows from $W_r < 0$, the second inequality from SCC along with the fact that

$$W(\bar{\theta}, \rho(\bar{\theta}), 0) = L(\bar{\theta}, \underline{r}) \leq W\left(\bar{\theta}, \underline{r}, \Psi\left(\frac{x^\#(\sigma) - \bar{\theta}}{\sigma}\right)\right)$$

and the last inequality from the monotonicity of the size of attack $\Psi\left(\frac{x^\#(\sigma) - \theta}{\sigma}\right)$ in θ along with the fact that $W_A \leq 0$. We thus conclude that, given the agents' strategy, no type of the policy maker has an incentive to raise the policy above $\underline{r}$.

To complete the proof, it then suffices to show that the above strategy profile for the agents can be supported by an appropriate system of beliefs. When $r = \underline{r}$, Bayes' rule imposes that, for any θ , $\mu(\theta|x, \underline{r}) = 1 - \Psi\left(\frac{x - \theta}{\sigma}\right)$. From the construction in the proof of Lemma 1, it is then immediate to see that, given these beliefs, attacking if and only if $x < x^\#(\sigma)$ is sequentially optimal for the

agents. Next, consider $r \in (\underline{r}, \rho(\bar{\theta})]$. Then let $\theta_r^* = \inf\{\theta \geq \underline{\theta} : W(\theta, r, 0) \geq L(\theta, \underline{r})\}$. Because for any $\theta \in [\theta_r^*, \bar{\theta}]$, $r \leq \rho(\theta)$, it then follows that, for any $\theta \in [\theta_r^*, \bar{\theta}]$, $Z(\theta, r) \geq Q(r)$. Then let $\mu(\cdot|x, r)$ be any beliefs that assign probability one to the event that $\theta \in [\theta_r^*, \bar{\theta}]$, irrespective of x . Because for any $\theta \in [\theta_r^*, \bar{\theta}]$, $D(\theta, r) = 1$, it then follows that these beliefs satisfy

$$\int_{\theta_r^*}^{\bar{\theta}} Z(\tilde{\theta}, r) d\mu(\tilde{\theta}|x, r) \geq Q(r) \text{ for all } x \quad (18)$$

which guarantees that any agent who expects all other agents to attack finds it optimal to do the same, irrespective of his signal. Finally, for any $r > \rho(\bar{\theta})$, let $\mu(\cdot|x, r)$ be any beliefs that assign probability one to $\theta > \underline{\theta}$ irrespective of x . Under such beliefs, any agent who expects no other agent to attack finds it optimal to refrain from attacking. *Q.E.D.*

Next, consider $s > \underline{r}$. We start by establishing the following preliminary result.

Lemma A1. *For any $\sigma > 0$ any $s \in (\underline{r}, r^\#(\sigma)]$, there exists at least one pair (θ_s^*, θ_s'') , with $\theta_s'' \geq \theta_s^*$ such that θ_s^* satisfies condition (8) and θ_s'' satisfies condition (10). Furthermore, $\theta_s'' > \theta_s^*$ if $s < r^\#(\sigma)$.*

Proof of Lemma A1. Fix $\sigma > 0$. By definition, $r^\#(\sigma)$ solves $W(\theta^\#(\sigma), r^\#(\sigma), 0) = L(\theta^\#(\sigma), \underline{r})$, where $\theta^\#(\sigma) \in (\underline{\theta}, \bar{\theta})$ is the unique threshold that corresponds to the pooling equilibria of parts (ii) and in Proposition (1). Because, for any θ , $W(\theta, \cdot, 0)$ is continuous and strictly decreasing in s and because, for any s , $W(\cdot, s, 0) - L(\cdot, \underline{r})$ is continuous and strictly increasing in θ , we have that, for any $s \in (\underline{r}, r^\#(\sigma)]$, $\{\theta \geq \underline{\theta} : W(\theta, s, 0) \geq L(\theta, \underline{r})\} \neq \emptyset$. Furthermore, $\theta_s^* \leq \theta^\#(\sigma)$ with strict inequality if and only if $s < r^\#(\sigma)$. Next note that $X(\theta, \theta; \sigma)$ coincides with the unique solution to

$$\int_{-\infty}^{\theta} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{x - \tilde{\theta}}{\sigma}\right) d\tilde{\theta} = Q(\underline{r})$$

From the results established in the proof of Lemma 1, one can then verify that $R(\theta, \underline{r}, B(\theta, \theta; \sigma)) \leq 0$ for all $\theta \leq \theta^\#(\sigma)$ with strict inequality for $\theta < \theta^\#(\sigma)$ and $R(\theta, \underline{r}, B(\theta, \theta; \sigma)) > 0$ for all $\theta > \theta^\#(\sigma)$.

Because $R(\cdot, \underline{r}, B(\theta_s^*, \cdot; \sigma))$ is continuous in θ and $R(\bar{\theta}, \underline{r}, B(\theta_s^*, \bar{\theta}; \sigma)) > 0$, this means that there always exists a $\theta_s'' \geq \theta_s^*$ (with strict inequality if $\theta_s^* < \theta^\#(\sigma)$, that is if $s < r^\#(\sigma)$) that satisfies condition (10). ■

Next, we show that, when SCC holds, then for any $\sigma > 0$, any $s > r^\#(\sigma)$, $\mathcal{E}(s; \sigma) = \emptyset$. To see this, note that, when $s > r^\#(\sigma)$, then either $\{\theta \geq \underline{\theta} : W(\theta, s, 0) \geq L(\theta, \underline{r})\} = \emptyset$ or, if the set is not empty, in which case its greatest lower bound is $\theta_s^* > \theta^\#(\sigma)$, then $R(\theta_s^*, \underline{r}, B(\theta_s^*, \theta_s^*; \sigma)) > 0$. Next observe that, because the function

$$B(\theta_s^*, \theta; \sigma) \equiv \Psi\left(\frac{X(\theta_s^*, \theta; \sigma) - \theta}{\sigma}\right)$$

is decreasing in θ over $[\theta_s^*, +\infty)$ (to see this, use the representation given in (28) in the proof of Proposition 9 below), then $R(\theta, \underline{r}, B(\theta_s^*, \theta; \sigma)) > 0$ for all $\theta \geq \theta_s^*$. This means that there is no $\theta_s'' \geq \theta_s^*$ that satisfies condition (10). To conclude that $\mathcal{E}(s; \sigma) = \emptyset$, it then suffices to show that there is also no type $\theta_s'' \geq \theta_s^*$ that satisfies condition (9). To see this, note that, for any $\theta \geq \theta_s^*$,

$$\begin{aligned} G(\theta; \theta_s^*, \sigma) &\equiv W(\theta, s, 0) - W(\theta, \underline{r}, B(\theta_s^*, \theta; \sigma)) \\ &\leq W(\theta, s, 0) - W(\theta, \underline{r}, B(\theta_s^*, \theta_s^*; \sigma)) \\ &< 0 \end{aligned}$$

where the first inequality follows again from the fact that $B(\theta_s^*, \theta; \sigma)$ is decreasing in θ , while the second inequality follows from SCC. We conclude that condition (9) admits no solution for $s > r^\#(\sigma)$. From Proposition 1, this means that $\mathcal{E}(s; \sigma) = \emptyset$.

Thus consider $s \in (\underline{r}, r^\#(\sigma)]$. That, for any $s \in (\underline{r}, r^\#(\sigma)]$, $\mathcal{E}(s; \sigma) \neq \emptyset$, as well as the existence of equilibria satisfying the properties of part (ii) in the proposition follows from Lemma A2 below.

Lemma A2. *Assume SCC holds. For any $\sigma > 0$, any $s \in (\underline{r}, r^\#(\sigma)]$, there exists an equilibrium in which $r(\theta) = s$ if $\theta \in (\theta_s^*, \theta_s^{**}(\sigma)]$ and $r(\theta) = \underline{r}$ otherwise. The equilibrium is sustained by the following strategy for the agents: (i) for $r = \underline{r}$, $a(x, r) = 1$ if and only if $x < x_s^* = X(\theta_s^*, \theta_s^{**}(\sigma); \sigma)$; for any $r \in (\underline{r}, s)$, $a(x, r) = 1$ irrespective of x ; for any $r \geq s$, $a(x, r) = 0$ irrespective of x .*

Proof of Lemma A2. From Lemma A1, for any $\sigma > 0$ any $s \in (\underline{r}, r^\#(\sigma)]$, there always exists a pair (θ_s^*, θ_s'') such that θ_s^* satisfies condition (8) and $\theta_s'' \geq \theta_s^*$ satisfies condition (10). Thus let $\theta_s^{**}(\sigma) = \sup\{\theta_s'' \geq \theta_s^* : \theta_s'' \text{ satisfies condition (9) or condition (10)}\}$. We then have that, for any $\theta \in (\theta_s^*, \theta_s^{**}(\sigma)]$,

$$\begin{aligned} &U(\theta, s, 0) - U\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta}{\sigma}\right)\right) \\ &= W(\theta, s, 0) - U\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta}{\sigma}\right)\right) \geq 0 \end{aligned}$$

To see this, note that the inequality trivially holds if $R\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta}{\sigma}\right)\right) \leq 0$, that is, if by not raising the policy, type θ experiences regime change, in which case $U\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta}{\sigma}\right)\right) = L(\theta, \underline{r}) < W(\theta, s, 0)$. Thus suppose that $R\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta}{\sigma}\right)\right) > 0$, which means that $\theta_s^{**}(\sigma)$

satisfies condition (9). In this case,

$$\begin{aligned}
& U(\theta, s, 0) - U\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta}{\sigma}\right)\right) \\
&= W(\theta, s, 0) - W\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta}{\sigma}\right)\right) \\
&\geq W(\theta, s, 0) - W\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta_s^{**}(\sigma)}{\sigma}\right)\right) \\
&\geq W(\theta_s^{**}(\sigma), s, 0) - W\left(\theta_s^{**}(\sigma), \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta_s^{**}(\sigma)}{\sigma}\right)\right) \\
&= 0
\end{aligned}$$

where the first inequality follows from the fact that W is non-increasing in A , the second inequality from SCC, and the last equality from the fact that $\theta_s^{**}(\sigma)$ satisfies condition (9). Similar arguments imply that for any $\theta > \theta_s^{**}(\sigma)$,

$$\begin{aligned}
& U(\theta, s, 0) - U\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta}{\sigma}\right)\right) \\
&= W(\theta, s, 0) - W\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta}{\sigma}\right)\right) \leq 0.
\end{aligned}$$

To see this, note that if $\theta_s^{**}(\sigma)$ satisfies condition (9), then the result follows from SCC along with the monotonicity of $W(\theta, \underline{r}, A)$ in A and the strict monotonicity of $\Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta}{\sigma}\right)$ in θ . If, instead,

$$\theta_s^{**}(\sigma) = \sup \left\{ \theta \geq \theta_s^* : R\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta; \sigma) - \theta}{\sigma}\right)\right) = 0 \right\}$$

then necessarily

$$W(\theta_s^{**}(\sigma), s, 0) - W\left(\theta_s^{**}(\sigma), \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta_s^{**}(\sigma)}{\sigma}\right)\right) \leq 0$$

But then, by the same arguments as above, for any $\theta > \theta_s^{**}(\sigma)$, $U(\theta, s, 0) - U\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta}{\sigma}\right)\right) \leq 0$.

We are now ready to establish the result in the lemma. Because $A(\theta, r) = 1 > A(\theta, \underline{r}) = \Psi\left(\frac{x_s^* - \theta}{\sigma}\right)$ for any $r \in (\underline{r}, s)$ and $A(\theta, r) = A(\theta, s) = 0$ for any $r \geq s$, the policy maker strictly prefers $\underline{r}$ to any $r \in (\underline{r}, s)$ and s to any $r > s$. Indeed, for any type $\theta \leq \bar{\theta}$, raising the policy to $r \in (\underline{r}, s)$ leads to regime change (with payoff $L(\theta, r) < L(\theta, \underline{r})$) whereas, for any type $\theta > \bar{\theta}$, raising the policy to $r \in (\underline{r}, s)$ yields a payoff $W(\theta, r, 1) < W(1, \underline{r}, A(\theta, \underline{r}))$. Likewise, raising the policy to $r > s$ leads to a lower payoff than raising the policy to s for any θ . Furthermore, $\underline{r}$ is dominant for any $\theta \leq \underline{\theta}$. For $\theta > \underline{\theta}$, on the other hand, the payoff from raising the policy to $r = s$ is $W(\theta, s, 0)$, while the payoff from leaving the policy at $r = \underline{r}$ is $U(\theta, \underline{r}, A(\theta, \underline{r}))$. From the definition of the thresholds θ_s^* and $\theta_s^{**}(\sigma)$ and the properties established above, we then have that raising the policy to $r = s$ is optimal if and only if $\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]$, which establishes the optimality of the policy maker's strategy.

Next, consider the agents. When $r = \underline{r}$, $D(\theta, r) = 1$ occurs if and only if $\theta \leq \hat{\theta}_s$, where $\hat{\theta}_s$ is the unique solution to

$$R(\hat{\theta}_s, \underline{r}, A(\hat{\theta}_s, \underline{r})) = 0.$$

When, instead, $r \in (\underline{r}, s)$, $D(\theta, r) = 1$ if and only if $\theta \leq \bar{\theta}$. Finally, for any $r \geq s$, $D(\theta, r) = 1$ if and only if $\theta \leq \underline{\theta}$. An agent thus finds it optimal to follow the equilibrium strategy if and only if his beliefs satisfy the following conditions:

$$\begin{aligned} \text{when } r = \underline{r}, \quad & \int_{-\infty}^{\hat{\theta}_s} Z(\tilde{\theta}, \underline{r}) d\mu(\tilde{\theta}|x, r) \geq Q(\underline{r}) \text{ if } x < x_s^* \\ & \text{and } \int_{-\infty}^{\hat{\theta}_s} Z(\tilde{\theta}, \underline{r}) d\mu(\tilde{\theta}|x, r) \leq Q(\underline{r}) \text{ if } x \geq x_s^*; \end{aligned} \quad (19)$$

$$\text{when } r \in (\underline{r}, s), \quad \int_{-\infty}^{\bar{\theta}} Z(\tilde{\theta}, \underline{r}) d\mu(\tilde{\theta}|x, r) \geq Q(r) \text{ for all } x; \quad (20)$$

$$\text{when } r \geq s, \quad \int_{-\infty}^{\underline{\theta}} Z(\tilde{\theta}, \underline{r}) d\mu(\tilde{\theta}|x, r) \leq Q(r) \text{ for all } x. \quad (21)$$

Beliefs are pinned down by Bayes' rule when either $r = \underline{r}$ or $r = s$. In the first case ($r = \underline{r}$), for any $\theta \leq \theta_s^*$

$$\mu(\theta|x, \underline{r}) = \frac{1 - \Psi\left(\frac{x - \theta}{\sigma}\right)}{1 - \Psi\left(\frac{x - \theta_s^*}{\sigma}\right) + \Psi\left(\frac{x - \theta_s^{**}(\sigma)}{\sigma}\right)},$$

while for any $\theta \in (\theta_s^*, \hat{\theta}_s)$,

$$\mu(\theta|x, \underline{r}) = \frac{1 - \Psi\left(\frac{x - \theta_s^*}{\sigma}\right)}{1 - \Psi\left(\frac{x - \theta_s^*}{\sigma}\right) + \Psi\left(\frac{x - \theta_s^{**}(\sigma)}{\sigma}\right)}.$$

That these beliefs satisfy condition (19) follows from the uniqueness of the threshold $X(\theta_s^*, \theta_s^{**}(\sigma); \sigma)$ which implies that the expected payoff from attacking changes sign only once at $x = x_s^* = X(\theta_s^*, \theta_s^{**}(\sigma); \sigma)$. In the second case ($r = s$), $\mu(\underline{\theta}|x, s) = 0$, in which case condition (21) is clearly satisfied. Finally, whenever $r \notin \{\underline{r}, s\}$, there exist an arbitrarily large set of out-of-equilibrium beliefs that satisfy (20) and (21) — see the construction in the proof of Lemma A0 for an example.

Combining the optimality of the agents' strategies with the optimality of the policy maker's strategy gives the result. *Q.E.D.*

Proof of Proposition 3. Most of the results in the proposition follow from Propositions 1 and 2 along with properties (1)-(2) below. To save on notation, hereafter we let $\theta_{\underline{r}}^*(\sigma) \equiv \theta_{\underline{r}}^{**}(\sigma) \equiv \theta^\#(\sigma)$ denote the unique thresholds corresponding to any of the pooling equilibria of $\mathcal{E}(\underline{r}; \sigma)$.

Property 1. For any $\sigma > 0$ any $s', s'' \in (\underline{r}, r^\#(\sigma)]$, $s'' > s'$ implies that $\Delta_{s''}(\sigma) \leq \Delta_{s'}(\sigma)$.

To see this, note that, for any $s \in (\underline{r}, r^\#(\sigma)]$,

$$\Delta_s(\sigma) \equiv \theta_s^{**}(\sigma) - \theta_s^* = \begin{cases} \hat{\Delta}_s(\sigma) & \text{if } \hat{G}(\Delta; s, \sigma) < 0 \text{ all } \Delta \geq \hat{\Delta}_s(\sigma) \\ \sup\{\Delta \geq \hat{\Delta}_s(\sigma) : \hat{G}(\Delta; s, \sigma) = 0\} & \text{otherwise} \end{cases}$$

where $\hat{\Delta}_s(\sigma)$ is the unique⁴⁰ $\Delta \geq 0$ that solves $R\left(\theta_s^* + \Delta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^* + \Delta; \sigma) - (\theta_s^* + \Delta)}{\sigma}\right)\right) = 0$ and where

$$\hat{G}(\Delta; s, \sigma) \equiv G(\theta_s^* + \Delta; \theta_s^*, \sigma) \equiv W(\theta_s^* + \Delta, s, 0) - W\left(\theta_s^* + \Delta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^* + \Delta; \sigma) - (\theta_s^* + \Delta)}{\sigma}\right)\right).$$

Now, for any $s \in (\underline{r}, r^\#(\sigma)]$, $\Delta \geq 0$, and $\sigma > 0$, let $\hat{B}(\Delta; s, \sigma) \equiv \Psi\left(\frac{X(\theta_s^*, \theta_s^* + \Delta; \sigma) - (\theta_s^* + \Delta)}{\sigma}\right)$ and note that $\hat{B}(\Delta; s, \sigma)$ is implicitly defined by

$$\int_{-\infty}^{\theta_s^*} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{X(\theta_s^*, \theta_s^* + \Delta; \sigma) - \tilde{\theta}}{\sigma}\right) d\tilde{\theta} - Q(\underline{r}) \left[1 - \Psi\left(\frac{X(\theta_s^*, \theta_s^* + \Delta; \sigma) - \theta_s^*}{\sigma}\right) + \hat{B}(\Delta; s, \sigma)\right] = 0.$$

Integrating by parts and using the definition of $\hat{B}(\Delta; s, \sigma)$ we then have that $\hat{B}(\Delta; s, \sigma)$ is implicitly defined by

$$\begin{aligned} & (Z(\theta_s^*, \underline{r}) - Q(\underline{r})) \left(1 - \Psi\left(\Psi^{-1}(\hat{B}) + \frac{\Delta}{\sigma}\right)\right) \\ & - \int_{-\infty}^{\theta_s^*} \frac{\partial Z(\tilde{\theta}, \underline{r})}{\partial \tilde{\theta}} \left(1 - \Psi\left(\Psi^{-1}(\hat{B}) + \frac{\theta_s^* + \Delta - \tilde{\theta}}{\sigma}\right)\right) d\tilde{\theta} = Q(\underline{r}) \hat{B} \end{aligned} \quad (22)$$

from which we obtain that, for any $s \in (\underline{r}, r^\#(\sigma))$, $\Delta \geq 0$, and $\sigma > 0$

$$\begin{aligned} \frac{\partial \hat{B}(\Delta; s, \sigma)}{\partial s} &= - \frac{\int_{-\infty}^{\theta_s^*} \frac{\partial Z(\tilde{\theta}, \underline{r})}{\partial \tilde{\theta}} \frac{1}{\sigma} \psi\left(\Psi^{-1}(\hat{B}) + \frac{\theta_s^* + \Delta - \tilde{\theta}}{\sigma}\right) \frac{d\theta_s^*}{ds} d\tilde{\theta}}{\left[\begin{aligned} & - (Z(\theta_s^*, \underline{r}) - Q(\underline{r})) \psi\left(\Psi^{-1}(\hat{B}) + \frac{\Delta}{\sigma}\right) \frac{d\Psi^{-1}(\hat{B})}{dx} \\ & + \int_{-\infty}^{\theta_s^*} \frac{\partial Z(\tilde{\theta}, \underline{r})}{\partial \tilde{\theta}} \psi\left(\Psi^{-1}(\hat{B}) + \frac{\theta_s^* + \Delta - \tilde{\theta}}{\sigma}\right) \frac{d\Psi^{-1}(\hat{B})}{dx} d\tilde{\theta} - Q(\underline{r}) \end{aligned} \right]} \\ &\leq 0. \end{aligned}$$

The result that $\hat{B}(\Delta; \cdot, \sigma) \equiv \Psi\left(\frac{X(\theta_s^*, \theta_s^* + \Delta; \sigma) - (\theta_s^* + \Delta)}{\sigma}\right)$ is non-increasing in s along with the fact that $R(\theta, \underline{r}, A)$ is increasing in θ and decreasing in A in turn imply that $\hat{\Delta}_s(\sigma)$ is non-increasing in s over $(\underline{r}, r^\#(\sigma)]$.

Next, consider the function $\hat{G}(\Delta; s, \sigma)$ and observe that, for any $\sigma > 0$, $\hat{G}(\cdot; \cdot, \sigma)$ is continuous in (s, Δ) . Now fix $s \in (\underline{r}, r^\#(\sigma))$ and $\sigma > 0$ and suppose that $\hat{G}(\Delta; s, \sigma) < 0$ for all $\Delta \geq \hat{\Delta}_s(\sigma)$, in which case $\Delta_s(\sigma) = \hat{\Delta}_s(\sigma)$. Because $\hat{G}(\Delta; s, \sigma)$ is continuous in (s, Δ) , there exists a $\delta > 0$ such that, for any $\hat{s} \in (s - \delta, s + \delta)$, any $\Delta \geq \hat{\Delta}_s(\sigma) - \delta$, $\hat{G}(\Delta; \hat{s}, \sigma) < 0$. Together with the monotonicity of $\hat{\Delta}_s(\sigma)$ in s , this means that $\partial \Delta_s(\sigma) / \partial s \leq 0$.

Finally, consider the case where, given $s \in (\underline{r}, r^\#(\sigma))$ and $\sigma > 0$, $\{\Delta \geq \hat{\Delta}_s(\sigma) : \hat{G}(\Delta; s, \sigma) = 0\} \neq \emptyset$ in which case $\Delta_s(\sigma) = \sup\{\Delta \geq \hat{\Delta}_s(\sigma) : \hat{G}(\Delta; s, \sigma) = 0\}$. The property that $\lim_{\Delta \rightarrow +\infty} \hat{G}(\Delta; s, \sigma) < 0$ implies that $\hat{G}(\cdot; s, \sigma)$ must be locally strictly decreasing in Δ at $\Delta = \Delta_s(\sigma)$; that is,

$$\frac{\partial \hat{G}(\Delta_s(\sigma); s, \sigma)}{\partial \Delta} < 0. \quad (23)$$

⁴⁰That such a solution exists follows from the fact that $s \leq r^\#(\sigma)$ along with the fact that $R\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta; \sigma) - \theta}{\sigma}\right)\right)$ is strictly increasing in θ , with $\lim_{\theta \rightarrow \theta_s^*} R\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta; \sigma) - \theta}{\sigma}\right)\right) < 0 < \lim_{\theta \rightarrow \bar{\theta}} R\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta; \sigma) - \theta}{\sigma}\right)\right)$.

Furthermore,

$$\begin{aligned} \frac{\partial \hat{G}(\Delta_s(\sigma); s, \sigma)}{\partial s} &= \left[W_\theta(\theta_s^* + \Delta_s(\sigma), s, 0) - W_\theta(\theta_s^* + \Delta_s(\sigma), \underline{r}, \hat{B}(\Delta_s(\sigma); \theta_s^*, \sigma)) \right] \times \frac{d\theta_s^*}{ds} \\ &\quad + W_r(\theta_s^* + \Delta_s(\sigma), s, 0) - W_A(\theta_s^* + \Delta_s(\sigma), \underline{r}, \hat{B}(\Delta_s(\sigma); s, \sigma)) \frac{\partial \hat{B}(\Delta_s(\sigma); s, \sigma)}{\partial s} \\ &\leq 0 \end{aligned} \quad (24)$$

That the second and third terms of (24) are nonpositive follows from the fact that W is decreasing in r and in A along with the fact that $\frac{\partial \hat{B}(\Delta_s(\sigma); s, \sigma)}{\partial s} \leq 0$ as shown above. That the first term of (24) is also non-positive follows from the fact that $W(\theta_s^* + \Delta_s(\sigma), s, 0) - W(\theta_s^* + \Delta_s(\sigma), \underline{r}, \hat{B}(\Delta_s(\sigma); \theta_s^*, \sigma)) = 0$. The SCC along with the fact that θ_s^* is locally nondecreasing then imply that the first term of (24) is nonpositive. From the Implicit Function Theorem, we then have that (23) together with (24) imply that $\partial \Delta_s(\sigma) / \partial s \leq 0$.

Combining the results above, then implies that, for any $\sigma > 0$ any $s', s'' \in (\underline{r}, r^\#(\sigma)]$, $s'' > s'$ implies that $\Delta_{s''}(\sigma) \leq \Delta_{s'}(\sigma)$.

Property 2. For any $\sigma, \sigma' > 0$, any $s \in (\underline{r}, \min\{r^\#(\sigma), r^\#(\sigma')\})$, $\sigma' > \sigma$ implies $\Delta_s(\sigma') \geq \Delta_s(\sigma)$.

This follows from the proof of part (iii) of Proposition 9 below, where it is shown that $\Delta_s(\sigma)$ is increasing in σ .

Given the aforementioned properties, the results in the proposition can be established as follows. First, note that, for any $\sigma > 0$, any $F \in \mathcal{F}(\sigma)$, any θ , any $r > \underline{r}$,

$$\begin{aligned} P(r, \theta; F, \sigma) &= \begin{cases} \int_{s \in [r, r^\#(\sigma)]} I_{\{\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}} dF(s) & \text{if } r \in (\underline{r}, r^\#(\sigma)] \\ 0 & \text{if } r > r^\#(\sigma) \end{cases} \\ \Delta(r; F, \sigma) &= \begin{cases} \int_{s \in [r, r^\#(\sigma)]} \Delta_s(\sigma) dF(s) & \text{if } r \in (\underline{r}, r^\#(\sigma)] \\ 0 & \text{if } r > r^\#(\sigma). \end{cases} \end{aligned}$$

Part (i). Fix $\sigma > 0$ and $F \in \mathcal{F}(\sigma)$, take any $r \in (\underline{r}, r^\#(\sigma)]$ and let $\theta^\circ(r; F, \sigma) = \theta_r^*$ and $\theta^{\circ\circ}(r; F, \sigma) = \sup\{\theta_s^{**}(\sigma) : s \in [r, r^\#(\sigma)]\}$. From Proposition 1, we then have that $P(r, \theta; F, \sigma) > 0$ only if $\theta \in [\theta^\circ(r; F, \sigma), \theta^{\circ\circ}(r; F, \sigma)]$.

Part (ii). Again, fix $\sigma > 0$ and take any $F \in \mathcal{F}(\sigma)$. That $D(\theta; F, \sigma)$ is non-increasing in θ , with $D(\theta; F, \sigma) = 1$ for $\theta < \underline{\theta}$ and $D(\theta; F, \sigma) = 0$ for $\theta > \theta^\#(\sigma)$ follows directly from Proposition 1.

Part (iii). Fix $\sigma > 0$ and take any pair $F, F' \in \mathcal{F}(\sigma)$. For any $\theta \in (\underline{\theta}, \theta^\#(\sigma))$,

$$D(\theta; F, \sigma) = F(\underline{r}) + 1 - F(\rho(\theta))$$

and similarly for F' . Because $F'(s) = F(s)$ for $s \in \{\underline{r}, r^\#(\sigma)\}$ and $F'(s) < F(s)$ for all $s \in (\underline{r}, r^\#(\sigma))$, then $D(\theta; F', \sigma) > D(\theta; F, \sigma)$, unless $D(\theta; F, \sigma) = 1$.

Next, consider the probability of intervention. From Property 1 above, $\Delta_s(\sigma)$ is a positive, (weakly) decreasing, and differentiable function of s , for any $s \in (\underline{r}, r^\#(\sigma))$. Now fix $r \in (\underline{r}, r^\#(\sigma))$. That $F(r) = F'(r)$ along with $F'(s) < F(s)$ for all $s \in (r, r^\#(\sigma))$ then imply that

$$\begin{aligned} \Delta(r; F, \sigma) - \Delta(r, F', \sigma) &= \int_{s \in [\underline{r}, r^\#(\sigma)]} \Delta_s(\sigma) dF(s) - \int_{s \in [\underline{r}, r^\#(\sigma)]} \Delta_s(\sigma) dF'(s) \\ &= -\Delta_r(\sigma)[F(r) - F'(r)] + \\ &\quad - \int_{s \in [r, r^\#(\sigma)]} \frac{d\Delta_s(\sigma)}{ds} [F(s) - F'(s)] ds \\ &\geq 0. \end{aligned}$$

Part (iv). Take any c.d.f. F with support $\text{Supp}[F]$ such that $F(s) = 0$ for all $s \leq r_1$ and $F(s) = 1$ for all $s \geq r_2$ for some $r_1, r_2 \in \mathbb{R}$ with $\underline{r} < r_1 < r_2 < \lim_{\sigma \rightarrow 0^+} r^\#(\sigma)$. Note that this implies that there exists $\bar{\sigma} > 0$, such that, for any $\sigma < \bar{\sigma}$, $\text{Supp}[F] \subset [\underline{r}, r^\#(\sigma)]$. That, for any $r > \underline{r}$, $\lim_{\sigma \rightarrow 0^+} \Delta(r, F, \sigma) = 0$ follows from the Dominated Convergence Theorem. To see this, let $H : \mathbb{R} \rightarrow \mathbb{R}$ be the function whose domain is $\text{Supp}[F]$ and that is given by

$$H(s) = \sup\{\theta : U(\theta, s, 0) - U(\theta, \underline{r}, 1) \geq 0\} - \theta_s^*$$

for all $s \in \text{Supp}[F]$. It is immediate that $H(s) \geq 0$, that $H(\cdot)$ is decreasing and that $\int H(s) dF(s) < \infty$. Furthermore, for any $\sigma < \bar{\sigma}$ and any $s \in \text{Supp}[F]$, $\Delta_s(\sigma) \leq H(s)$. These properties, together with the result in Proposition 9 in Appendix A, then imply that

$$\lim_{\sigma \rightarrow 0^+} \Delta(r, F, \sigma) = \lim_{\sigma \rightarrow 0^+} \int_{s \geq r} \Delta_s(\sigma) dF(s) = \int_{s \geq r} \lim_{\sigma \rightarrow 0^+} \Delta_s(\sigma) dF(s) = 0.$$

That, for any θ , any $\sigma, \sigma' > 0$, any $F \in \mathcal{F}(\sigma) \cap \mathcal{F}(\sigma')$, $D(\theta; F, \sigma) = D(\theta; F, \sigma')$ follows from Proposition 1.

Lastly, that $\sigma' > \sigma > 0$ implies $\Delta(r, F, \sigma') \geq \Delta(r, F, \sigma)$ for all $r \in (\underline{r}, \min\{r^\#(\sigma), r^\#(\sigma')\})$ all $F \in \mathcal{F}(\sigma) \cap \mathcal{F}(\sigma')$, follows from Property 2 above. *Q.E.D.*

Proof of Proposition 4. Part (i). To see that $\underline{D}(\theta_1, \theta_2; \sigma)$ and $\bar{D}(\theta_1, \theta_2; \sigma)$ are both non-increasing in (θ_1, θ_2) , in the weak-order sense, consider any pair $(\theta_1, \theta_2), (\theta'_1, \theta'_2)$ such that $\theta_1 \leq \theta'_1$ and $\theta_2 \leq \theta'_2$. Clearly, the distribution of $\tilde{\theta}$ conditional on the event that $\theta \in (\theta'_1, \theta'_2)$ first-order-stochastically dominates the distribution of $\tilde{\theta}$ conditional on the event that $\theta \in (\theta_1, \theta_2)$. Along with the fact that, for any $F \in \mathcal{F}(\sigma)$, $D(\cdot; F, \sigma)$ is non-increasing in θ , this means that, for any $F \in \mathcal{F}(\sigma)$, $D(\theta_1, \theta_2; F, \sigma) \geq D(\theta'_1, \theta'_2; F, \sigma)$. Standard envelope arguments, then imply that the same monotonicities apply to $\underline{D}(\theta_1, \theta_2; \sigma)$ and $\bar{D}(\theta_1, \theta_2; \sigma)$. The result for $\underline{P}(r, \theta_1, \theta_2; \sigma)$ and $\bar{P}(r, \theta_1, \theta_2; \sigma)$ follows directly from Proposition 3.

Part (ii). That $\underline{P}(r, \theta_1, \theta_2; \sigma)$ is independent of σ is immediate. That also $\underline{D}(\theta_1, \theta_2; \sigma)$ is independent of σ follows from the fact that, for any $\sigma > 0$, $\lim_{s \rightarrow \underline{r}^+} \theta_s^* = \underline{\theta}$ along with the fact that

$D_s(\theta; \sigma) = 0$ for any $\theta > \theta_s^*$ in any equilibrium in $\mathcal{E}(s; \sigma)$. That, for any (θ_1, θ_2) , any $\sigma, \sigma' > 0$, $\theta_2 < \min\{\theta^\#(\sigma'), \theta^\#(\sigma)\}$ or $\theta_1 > \max\{\theta^\#(\sigma'), \theta^\#(\sigma)\}$ imply $\bar{D}(\theta_1, \theta_2; \sigma) = \bar{D}(\theta_1, \theta_2; \sigma')$ follows from the fact that, for any $\sigma > 0$, any θ , $D_s(\theta; \sigma) \leq D_{\underline{r}}(\theta; \sigma)$ together with the fact that $D_{\underline{r}}(\theta; \sigma) = 1$ for all $\theta \leq \theta^\#(\sigma)$ while $D_{\underline{r}}(\theta; \sigma) = 0$ for all $\theta > \theta^\#(\sigma)$.

Next, consider the claim that $\lim_{\sigma \rightarrow 0^+} \bar{P}(r, \theta_1, \theta_2; \sigma) = 0$ for any $r > \underline{r}$ and any $\theta_1, \theta_2 \in \mathbb{R}$. Clearly, the result is true if $r > r^\#(0^+) \equiv \lim_{\sigma \rightarrow 0^+} r^\#(\sigma)$. Thus take $r \in (\underline{r}, r^\#(0^+)]$. Let $\mathcal{F}$ denote an arbitrary set of c.d.f.s F with support $\text{Supp}[F] \subset [\underline{r}, \rho(\bar{\theta})]$ with the following properties: (i) $\mathcal{F}(\sigma) \subset \mathcal{F}$ for any $\sigma > 0$; (ii) $\mathcal{F}$ is compact with respect to the metric $d(\cdot)$ defined, for any pair $F_1, F_2 \in \mathcal{F}$, by $d(F_1, F_2) \equiv \sup\{|F_1(A) - F_2(A)| : A \in \Sigma\}$, where Σ is the Borel sigma algebra associated with the interval $[\underline{r}, \rho(\bar{\theta})]$. For any $\sigma > 0$, any c.d.f. $F \in \mathcal{F}$, any $\theta_1 < \theta_2$, and any $r > \underline{r}$, then let

$$\hat{P}(r, \theta_1, \theta_2; F, \sigma) \equiv \frac{1}{\theta_2 - \theta_1} \int_{\theta_1}^{\theta_2} \int_{s \in [\underline{r}, \rho(\bar{\theta})]} I_{\{\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}} dF(s) d\theta$$

with the convention that, given any $s \in [\underline{r}, \rho(\bar{\theta})]$, $I_{\{\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}} = 0$ if there is no $\theta_s'' \geq \theta_s^*$ that satisfies condition (9) or condition (10). Note that

$$\hat{P}(r, \theta_1, \theta_2; F, \sigma) = P(r, \theta_1, \theta_2; F, \sigma) \equiv \frac{1}{\theta_2 - \theta_1} \int_{\theta_1}^{\theta_2} P(r, \theta; F, \sigma) d\theta$$

if $F \in \mathcal{F}(\sigma)$, that is, if, given $\sigma > 0$, one restricts F to have support $\text{Supp}[F] \subset [\underline{r}, r^\#(\sigma)]$.

By Proposition 9 in Appendix A, for any $s \leq r^\#(0^+)$, $\lim_{\sigma \rightarrow 0^+} \Delta_s(\sigma) = 0$. This implies that, for any $\theta \in [\theta_1, \theta_2]$, any $s \in [\underline{r}, \rho(\bar{\theta})]$ with $\theta_s^* \neq \theta$,

$$\lim_{\sigma \rightarrow 0^+} I_{\{\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}} = 0.$$

By the Lebesgue Dominated Convergence Theorem, we then have that, for any $\theta \in [\theta_1, \theta_2]$

$$\lim_{\sigma \rightarrow 0^+} \int_{s \in [\underline{r}, \rho(\bar{\theta})]} I_{\{\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}} dF(s) = \int_{s \in [\underline{r}, \rho(\bar{\theta})]} \lim_{\sigma \rightarrow 0^+} I_{\{\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}} dF(s)$$

where the last integral is equal to zero, unless F has a mass point at an s such that $\theta_s^* = \theta$. It follows that

$$\lim_{\sigma \rightarrow 0^+} \int_{\theta_1}^{\theta_2} \left\{ \int_{s \in [\underline{r}, \rho(\bar{\theta})]} I_{\{\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}} dF(s) \right\} d\theta = \int_{\theta_1}^{\theta_2} \left\{ \lim_{\sigma \rightarrow 0^+} \int_{s \in [\underline{r}, \rho(\bar{\theta})]} I_{\{\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}} dF(s) \right\} d\theta = 0, \quad (25)$$

where the first equality is again by the Dominated Convergence Theorem, while the second equality follows from the following property. Given any c.d.f. F with support $\text{Supp}[F] \subset [\underline{r}, \rho(\bar{\theta})]$, there does not exist a Lebesgue positive-measure set $E \subset [\theta_1, \theta_2]$ such that, for any $\theta \in E$, $\theta_s^* = \theta$ with strictly positive probability. Formally, the set

$$S \equiv \left\{ s \in [\underline{r}, \rho(\bar{\theta})] : F(s) > \lim_{x \rightarrow s^-} F(x) \right\}$$

has zero Lebesgue measure, which in turn implies that the set

$$\Theta^+ \equiv \left\{ \theta \in [\theta_1, \theta_2] : \theta_s^* = \theta \text{ with } F(s) > \lim_{x \rightarrow s^-} F(x) \right\}$$

also has zero Lebesgue measure. This means that the set of points $\theta \in [\theta_1, \theta_2]$ such that

$$\lim_{\sigma \rightarrow 0^+} \int_{s \in [r, \rho(\bar{\theta})]} I_{\{\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}} dF(s) > 0$$

has zero Lebesgue measure, which implies the result in (25).

Having established that, for any c.d.f. $F \in \mathcal{F}$, $\lim_{\sigma \rightarrow 0^+} \hat{P}(r, \theta_1, \theta_2; F, \sigma) = 0$, we now show that this property also implies that

$$\lim_{\sigma \rightarrow 0^+} \left\{ \sup_{F \in \mathcal{F}(\sigma)} P(r, \theta_1, \theta_2; F, \sigma) \right\} = 0.$$

First note that, by definition, for any $\sigma > 0$,

$$\sup_{F \in \mathcal{F}(\sigma)} P(r, \theta_1, \theta_2; F, \sigma) \leq \sup_{F \in \mathcal{F}} \hat{P}(r, \theta_1, \theta_2; F, \sigma)$$

which implies that

$$\lim_{\sigma \rightarrow 0^+} \left\{ \sup_{F \in \mathcal{F}(\sigma)} P(r, \theta_1, \theta_2; F, \sigma) \right\} \leq \lim_{\sigma \rightarrow 0^+} \left\{ \sup_{F \in \mathcal{F}} \hat{P}(r, \theta_1, \theta_2; F, \sigma) \right\}. \quad (26)$$

To establish the result, it thus suffices to show that the right hand side of (26) is zero. This is established as follows. First, recall that, by assumption, $\mathcal{F}$ is compact in the metric $d(F_1, F_2) \equiv \sup \{|F_1(A) - F_2(A)| : A \in \Sigma\}$. Because it is metric, then $\mathcal{F}$ is also Hausdorff. Next, note that, for any $\bar{\sigma} > 0$, the function family $\left\{ \hat{P}(r, \theta_1, \theta_2; \cdot, \sigma) \right\}_{\sigma \in (0, \bar{\sigma}]}$ with domain $\mathcal{F}$ and range in $[0, 1]$ is uniform equicontinuous in the metric $d(\cdot)$ defined above. This means that for any $\varepsilon > 0$, there exists a $\delta > 0$ (which may depend on ε only) such that for any $\sigma \in (0, \bar{\sigma}]$ (i.e., for any family member $\hat{P}(r, \theta_1, \theta_2; \cdot, \sigma)$), any $F_1, F_2 \in \mathcal{F}$ such that $d(F_1, F_2) < \delta$,

$$|\hat{P}(r, \theta_1, \theta_2; F_1, \sigma) - \hat{P}(r, \theta_1, \theta_2; F_2, \sigma)| < \varepsilon.$$

To see that this is true, note that

$$\begin{aligned} & |\hat{P}(r, \theta_1, \theta_2; F_1, \sigma) - \hat{P}(r, \theta_1, \theta_2; F_2, \sigma)| \\ &= \frac{1}{\theta_2 - \theta_1} \left| \int_{\theta_1}^{\theta_2} \left[\int_{s \in [r, \rho(\bar{\theta})]} I_{\{\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}} dF_1(s) - \int_{s \in [r, \rho(\bar{\theta})]} I_{\{\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}} dF_2(s) \right] d\theta \right| \\ &= \frac{1}{\theta_2 - \theta_1} \left| \int_{\theta_1}^{\theta_2} [F_1(A(\theta, \sigma)) - F_2(A(\theta, \sigma))] d\theta \right| \end{aligned}$$

where $A(\theta, \sigma) \equiv \{s \in [r, \rho(\bar{\theta})] : \theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}$. It is then easy to see that, for any $\sigma \in (0, \bar{\sigma}]$, the result follows by letting $\delta = \varepsilon$. It is also easy to see that the function family $\left\{ \hat{P}(r, \theta_1, \theta_2; \cdot, \sigma) \right\}_{\sigma \in (0, \bar{\sigma}]}$ is uniformly bounded, i.e., that there exists $M > 0$ such that $|\hat{P}(r, \theta_1, \theta_2; F, \sigma)| < M$ all $F \in \mathcal{F}$, all

$\sigma \in (0, \bar{\sigma}]$. From the Ascoli-Arzelà Theorem, any sequence of equicontinuous, uniformly bounded, functions defined on a compact Hausdorff space has a uniformly convergent subsequence (i.e., a subsequence that is convergent in the sup-norm). This implies $\lim_{\sigma \rightarrow 0^+} \left\{ \sup_{F \in \mathcal{F}} \hat{P}(r, \theta_1, \theta_2; F, \sigma) \right\} = 0$.

Lastly, we show that when $Z(\theta, \underline{r}) = z > \underline{r}$ for all θ , then the bound $\bar{D}$ is independent of σ whereas the bound $\bar{P}(r, \theta_1, \theta_2; \sigma)$ is a nondecreasing function of σ . The first property follows directly from the fact that, in this case, $\theta^\#(\sigma)$ and hence $r^\#(\sigma)$ are independent of $\sigma > 0$, together with the fact that, for any $\theta > \theta^\#$, $D_s(\theta; \sigma) = 0$ for all s , while for any $\theta \leq \theta^\#$ $D_{\underline{r}}(\theta, \sigma) = 1$ for all $\sigma > 0$. The second property follows from the fact that, given any distribution F with support⁴¹ $Supp[F] \subset [\underline{r}, r^\#]$,

$$P(r, \theta_1, \theta_2; F, \sigma) = \frac{1}{\theta_2 - \theta_1} \int_{\theta_1}^{\theta_2} \int_{s \in [r, r^\#]} I_{\{\theta \in [\theta_s^*, \theta_s^{**}(\sigma)]\}} dF(s) d\theta$$

That $P(r, \theta_1, \theta_2; F, \sigma)$ is increasing in σ then follows from the fact that $\theta_s^{**}(\sigma)$ is increasing in σ . By the envelope theorem, that each $P(r, \theta_1, \theta_2; F, \sigma)$ is weakly increasing in σ , then implies that the upper bound $\bar{P}(r, \theta_1, \theta_2; \sigma)$ is also weakly increasing in σ , which establishes the result. *Q.E.D.*

Proof of Proposition 5. For any θ , let

$$\Delta \Pi_s(\theta; \sigma) \equiv W(\theta, s, 0) - U\left(\theta, \underline{r}, \Psi\left((x^\#(\sigma) - \theta)/\sigma\right)\right)$$

denote the difference between the payoff $W(\theta, s, 0)$ that type θ obtains by raising the policy to $r = s$ in case regime change does not occur and no agent attacks in the game in which policy interventions are possible and the equilibrium payoff $U^o(\theta; \sigma) = U\left(\theta, \underline{r}, \Psi\left((x^\#(\sigma) - \theta)/\sigma\right)\right)$ that the same type obtains in the game in which the option to intervene is absent. The fact that $s < r^\#(\sigma)$ implies that, for any type $\theta \in (\theta_s^*, \theta^\#(\sigma)]$, $\Delta \Pi_s(\theta; \sigma) = W(\theta, s, 0) - L(\theta, \underline{r}) > 0$; indeed any such type, by raising the policy to $r = s$, can guarantee himself a payoff $W(\theta, s, 0)$ that is strictly higher than the payoff $U^o(\theta; \sigma) = L(\theta, \underline{r})$ that the same type would obtain absent the possibility to intervene. Furthermore, for any $\theta > \theta^\#(\sigma)$, $\Delta \Pi_s(\theta; \sigma) = W(\theta, s, 0) - W(\theta, \underline{r}, \Psi((x^\#(\sigma) - \theta)/\sigma))$. The fact that $\Delta \Pi_s(\theta; \sigma)$ is continuous over $(\theta^\#(\sigma), +\infty)$ together with SCC and the limit condition that $\lim_{\theta \rightarrow \infty} \rho(\theta) = \underline{r}$ (which, recall, is equivalent to $\lim_{\theta \rightarrow +\infty} W(\theta, s, 0) - W(\theta, \underline{r}, 1) < 0$), imply that there exists a unique $\theta_s^\dagger(\sigma) \geq \theta^\#(\sigma)$ such that $\Delta \Pi_s(\theta; \sigma) > 0$ for $\theta \in [\theta^\#(\sigma), \theta_s^\dagger(\sigma)]$ and $\Delta \Pi_s(\theta; \sigma) < 0$ for $\theta > \theta_s^\dagger(\sigma)$. Hence, no matter which particular equilibrium in $\mathcal{E}(s; \sigma)$ is played, any type $\theta \in (\theta_s^*, \theta_s^\dagger(\sigma)]$ is necessarily strictly better off with the option to intervene, whereas any type $\theta \leq \theta_s^*$ is just as well off.

Next note that any type $\theta > \theta_s^\dagger(\sigma)$ can be strictly worse off with the option to intervene only if the attack he expects when leaving the policy at $r = \underline{r}$ is strictly greater than the attack he would have

⁴¹Recall that in this case $\mathcal{F}(\sigma) = \mathcal{F}(\sigma')$ all $\sigma, \sigma' > 0$.

faced without the option to intervene, which is possible if and only if $X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) > x^\#(\sigma)$.⁴² However, σ small enough ensures that $X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) < x^\#(\sigma)$ and hence that the policy maker is always better off with the option to intervene, no matter his type. This follows from the fact that, when $\sigma \rightarrow 0^+$, $\theta_s^{**}(\sigma) \rightarrow \theta_s^*$ and $X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) \rightarrow \theta_s^*$ in the game with policy intervention,⁴³ whereas $x^\#(\sigma) \rightarrow \theta^\#(\sigma)$ in the game in which the option to intervene is absent. Together with the fact that $\theta_s^* < \theta^\#(\sigma)$ then gives the result. *Q.E.D.*

Proof of Proposition 6. First, consider the supremum of the equilibrium payoffs. For any $\theta > \underline{\theta}$, the highest *feasible* payoff is $W(\theta, \underline{r}, 0)$, the payoff enjoyed when regime change is avoided without facing any attack and without incurring any cost of policy intervention. This payoff can be approximated arbitrarily well in the game in which intervention is possible (it suffices to take any equilibrium in which s is sufficiently close to $\underline{r}$), which implies that $\bar{U}(\theta; \sigma) \geq U^o(\theta; \sigma)$ for all $\theta > \underline{\theta}$. That $\bar{U}(\theta; \sigma) > U^o(\theta; \sigma)$ for $\theta \in (\underline{\theta}, \theta^\#(\sigma)]$ follows directly from the fact that, without the option to intervene, these types experience regime change. That $\lim_{\theta \rightarrow +\infty} |\bar{U}(\theta; \sigma) - U^o(\theta; \sigma)| = 0$ follows from the fact, in the game without the option to intervene, $\lim_{\theta \rightarrow +\infty} A(\theta, \underline{r}) = 0$. This establishes the first part of the proposition.

Next consider the infimum of the equilibrium payoffs. As explained in the proof of Proposition 5, type θ can be worse off with the option to intervene only if there exists an $s \in (\underline{r}, r^\#(\sigma)]$ such that $X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) > x^\#(\sigma)$; recall that this means that there exists an equilibrium in which the attack that type θ faces when he does not raise the policy above $\underline{r}$ is larger than in the game without the option to intervene. Now suppose that such an s exists. From the result in Proposition 5, in this case any type $\theta \leq \theta_s^\dagger(\sigma)$ is still weakly better off (strictly for $\theta \in (\theta_s^*, \theta_s^\dagger(\sigma))$). However, by taking the equilibrium in $\mathcal{E}(s; \sigma)$ in which all agents attack when $r = \underline{r}$ if and only if $x < X(\theta_s^*, \theta_s^{**}(\sigma); \sigma)$ ensures that, for all $\theta > \theta_s^\dagger(\sigma)$, the payoff

$$\max \{W(\theta, s, 0); U(\theta, \underline{r}, \Psi((X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta)/\sigma))\}$$

that type θ obtains with the option to intervene is (weakly) lower than his payoff in the game in which the option to intervene is absent.⁴⁴ This means that there exists $\theta^\dagger(\sigma)$ such that $\underline{U}(\theta; \sigma) < U^o(\theta; \sigma)$

⁴²Recall from the analysis in Section 3, that (i) in any equilibrium in $\mathcal{E}(s; \sigma)$ no agent attacks when, after observing $r = \underline{r}$, he receives a signal $x > X(\theta_s^*, \theta_s^{**}(\sigma); \sigma)$, and (ii) that there always exists one equilibrium in $\mathcal{E}(s; \sigma)$ such that, after observing $r = \underline{r}$, each agent attacks if and only if he receives a signal $x < X(\theta_s^*, \theta_s^{**}(\sigma); \sigma)$.

⁴³The first property follows directly from Proposition 9 in Appendix A. To see that the second property also holds, note that, if this were not true, then the size of attack $A(\theta_s^{**}(\sigma), \underline{r})$ at $\theta = \theta_s^{**}(\sigma)$ would converge to either zero or one, making it impossible for type $\theta_s^{**}(\sigma)$ to be indifferent between raising the policy to $r = s$ and setting $r = \underline{r}$.

⁴⁴To see this, note that if $W(\theta, s, 0) > U(\theta, \underline{r}, \Psi((X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta)/\sigma))$ the result follows from the definition of $\theta_s^\dagger(\sigma)$. If, instead, the inequality is reversed, the result follows from the fact that, by not raising the policy, type θ faces an attack $\Psi((X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) - \theta)/\sigma)$ greater than the attack $\Psi((x^\#(\sigma) - \theta)/\sigma)$ he would have faced without the option to intervene.

only if $\theta > \theta^\dagger(\sigma)$. That $\theta^\dagger(\sigma) > \theta^\#(\sigma)$ is immediate given that there is no equilibrium in which a type $\theta \leq \theta^\#(\sigma)$ can be made worse off.

Finally, to see why, for any θ , the difference between $\underline{U}(\theta; \sigma)$ and $U^o(\theta; \sigma)$ vanishes as $\sigma \rightarrow 0^+$, note that, for any $\theta \leq \theta^\#(0^+) \equiv \lim_{\sigma \rightarrow 0^+} \theta^\#(\sigma)$, this difference is clearly zero because the lower bound is simply the payoff obtained in any equilibrium in which type θ is forced to experience regime change. For types $\theta > \theta^\#(0^+)$, on the other hand, the result follows from the fact that for any $s \in (\underline{r}, r^\#(0^+)]$, $\lim_{\sigma \rightarrow 0^+} \Delta_s(\sigma) = 0$ and

$$\lim_{\sigma \rightarrow 0^+} X(\theta_s^*, \theta_s^{**}(\sigma); \sigma) = \theta_s^* \leq \lim_{\sigma \rightarrow 0^+} x^\#(\sigma)$$

This implies that the lowest bound on the payoff for each $\theta > \theta^\#(0^+)$ is attained under any of the pooling equilibria, which is clearly the same payoff as in the game in which the option to intervene is absent. *Q.E.D.*

Proof of Proposition 7. For $\theta \leq \underline{\theta}$, it is dominant for the policy maker to set $\underline{r}$. Similarly, for $\theta > \bar{\theta}$, it is dominant for the agents not to attack, which makes it iteratively dominant for the policy maker to set $r = \underline{r}$. Finally, take any $\theta \in (\underline{\theta}, \bar{\theta}]$. Clearly, there is no subgame-perfect equilibrium in which $r(\theta) > \rho(\theta)$. On the other hand, for any $r \leq \rho(\theta)$, the assumption that $Z(\theta, r) \geq r$ implies that the continuation game among the agents is a coordination game with two Nash equilibria—one where nobody attacks and the status quo is maintained and another where everybody attacks and the status quo is abandoned. This implies that, for any $r' \leq \rho(\theta)$, there exists a subgame-perfect equilibrium in which the agents attack if and only if $r < r'$ and the policy maker raises the policy at $r(\theta) = r'$. This establishes part (i). Part (ii) follows from the properties above along with the fact that, for any $\theta \in (\underline{\theta}, \bar{\theta}]$ there always exists a subgame-perfect equilibrium where all agents attack if and only if $r \leq \rho(\theta)$ in which case the policy maker optimally chooses not to intervene and hence $D(\theta) = 1$. *Q.E.D.*

Proof of Proposition 8. The characterization of $\mathcal{G}(0)$ follows directly from Proposition 7. Thus consider $\mathcal{G}(\sigma)$ for $\sigma > 0$ and note that this set is given by

$$\mathcal{G}(\sigma) = \left\{ \begin{array}{l} (\theta, r): \text{ either } r = \underline{r} \text{ and } \theta \in \mathbb{R}, \\ \text{or } r \in (\underline{r}, \underline{\rho}^+] \text{ and } \theta \in (\theta_r^*, \theta_r^{**}], \text{ where } \theta_r^* = \underline{\theta}, \\ \text{or } r \in (\underline{\rho}^+, r^\#(\sigma)] \text{ and } \theta \in [\theta_r^*, \theta_r^{**}], \text{ where } \theta_r^* > \underline{\theta}. \end{array} \right\}$$

Next, pick an arbitrary $\bar{\varepsilon} \in (0, \rho(\bar{\theta}) - r^\#(0^+))$ if $\underline{\rho}^+ = \underline{r}$ and some $\bar{\varepsilon} \in (0, \underline{\rho}^+ - \underline{r})$ if $\underline{\rho}^+ > \underline{r}$, and for any $\varepsilon < \bar{\varepsilon}$ define the set $\mathcal{H}(\varepsilon)$ as follows: if $\underline{\rho}^+ = \underline{r}$, then

$$\mathcal{H}(\varepsilon) \equiv \left\{ \begin{array}{l} (\theta, r): \text{ either } r = \underline{r} \text{ and } \theta \in \mathbb{R}, \\ \text{or } r \in (\underline{r}, \underline{r} + \varepsilon) \text{ and } \theta > \underline{\theta} \\ \text{or } r \in [\underline{r} + \varepsilon, r^\#(0^+) + \varepsilon] \text{ and } \theta \in [\rho^{-1}(r), \rho^{-1}(r) + \varepsilon] \end{array} \right\};$$

whereas, if $\underline{\rho}^+ > \underline{r}$, then

$$\mathcal{H}(\varepsilon) \equiv \left\{ \begin{array}{l} (\theta, r): \text{ either } r = \underline{r} \text{ and } \theta \in \mathbb{R}, \\ \text{ or } r \in (\underline{r}, \underline{r} + \varepsilon) \text{ and } \theta > \underline{\theta} \\ \text{ or } r \in [\underline{r} + \varepsilon, \underline{\rho}^+] \text{ and } \theta \in (\underline{\theta}, \underline{\theta} + \varepsilon) \\ \text{ or } r \in [\underline{\rho}^+, r^\#(0^+) + \varepsilon] \text{ and } \theta \in [\rho^{-1}(r), \rho^{-1}(r) + \varepsilon] \end{array} \right\}.$$

Note that, in either case, the results in Proposition 9 in Appendix A guarantee that, for any $\varepsilon \in (0, \bar{\varepsilon})$, there exists $\bar{\sigma} > 0$ such that

$$\mathcal{G}(\sigma) \subset \mathcal{H}(\varepsilon)$$

for all $\sigma \in (0, \bar{\sigma})$. Finally, note that, for all $\sigma > 0$,

$$\mathcal{G}(\sigma) \supset \left\{ \begin{array}{l} (\theta, r): \text{ either } r = \underline{r} \text{ and } \theta \in \mathbb{R}, \\ \text{ or } r \in (\underline{\rho}^+, r^\#(\sigma)] \text{ and } \theta = \rho^{-1}(r) \end{array} \right\}$$

Combining the above two properties gives the result. *Q.E.D.*

Proof of Proposition 9. Fix $\varepsilon > 0$ and let

$$\theta^\#(0^+) \equiv \lim_{\sigma \rightarrow 0^+} \theta^\#(\sigma) \text{ and } r^\#(0^+) \equiv \lim_{\sigma \rightarrow 0^+} r^\#(\sigma) = \lim_{\sigma \rightarrow 0^+} \hat{r}(\theta^\#(\sigma))$$

where, for any $\theta \geq \underline{\theta}$, $\hat{r}(\theta)$ is implicitly defined by $W(\theta, \hat{r}(\theta), 0) = L(\theta, \underline{r})$.

Part (i). Clearly, irrespective of σ , $\mathcal{E}(s; \sigma) = \emptyset$ whenever $\{\theta \geq \underline{\theta} : W(\theta, s, 0) \geq L(\theta, \underline{r})\} = \emptyset$. Thus assume that this set is non-empty and denote by θ_s^* its largest lower bound. We establish the result using the properties below.

Property 1. *There exist $\delta > 0$ and $\hat{\sigma} > 0$ such that for any $\sigma < \hat{\sigma}$, $\theta_s^* > \theta^\#(0^+) + \delta$ whenever $s > r^\#(\sigma) + \varepsilon$.*

To see this, take $\varepsilon' < \varepsilon/2$. By continuity of $\hat{r}(\theta)$ and $\theta^\#(\sigma)$ and the fact that $r^\#(\sigma) = \hat{r}(\theta^\#(\sigma))$, there exists $\hat{\sigma} > 0$ such that, whenever $\sigma < \hat{\sigma}$, $|r^\#(\sigma) - r^\#(0^+)| < \varepsilon'$. Hence, for any $\sigma < \hat{\sigma}$, $s > r^\#(\sigma) + \varepsilon$ implies that $s > r^\#(0^+) - \varepsilon' + \varepsilon > r^\#(0^+) + \varepsilon'$. Furthermore, by the continuity of $\hat{r}(\theta)$, there exists $\delta > 0$ such that $|\hat{r}(\theta) - r^\#(0^+)| < \varepsilon'$ whenever $|\theta - \theta^\#(0^+)| < \delta$. By the monotonicity of $\hat{r}(\theta)$, we then have that $\hat{r}(\theta) < s$ for any $\theta < \theta^\#(0^+) + \delta$ which means that $\theta_s^* \geq \theta^\#(0^+) + \delta$.

Property 2. *Let $\hat{\sigma}$ be the threshold in property 1. There exist $0 < \sigma' < \hat{\sigma}$ and $\bar{\theta} > \underline{\theta}$ such that, for any $\sigma < \sigma'$, any $s > r^\#(\sigma) + \varepsilon$, any A , any $\theta > \bar{\theta}$, $R(\theta, \underline{r}, A) > 0$ and $W(\theta, s, 0) < W(\theta, \underline{r}, A)$.*

This follows from the limit condition that, for any $s > \underline{r}$, $\lim_{\theta \rightarrow +\infty} [W(\theta, s, 0) - W(\theta, \underline{r}, 1)] < 0$.

Property 3. *For any (θ^*, θ) , with $\theta \geq \theta^* \geq \underline{\theta}$, $B(\theta^*, \theta; \cdot)$ is increasing in σ with*

$$\lim_{\sigma \rightarrow 0^+} B(\theta^*, \theta; \sigma) = \begin{cases} 0 & \text{if } \theta > \theta^* \\ (Z(\theta^*) - Q(\underline{r})) / Z(\theta^*) & \text{if } \theta = \theta^* \end{cases}$$

Recall that $B(\theta^*, \theta; \sigma)$ is defined as

$$B(\theta^*, \theta; \sigma) \equiv \Psi \left(\frac{X(\theta^*, \theta; \sigma) - \theta}{\sigma} \right), \quad (27)$$

where $X(\theta^*, \theta; \sigma)$ is implicitly defined by condition (5). To simplify the exposition, let us momentarily use X and B as short-cuts for, respectively, $X(\theta^*, \theta; \sigma)$ and $B(\theta^*, \theta; \sigma)$. Condition (5) can then be restated as

$$\int_{-\infty}^{\theta^*} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi \left(\frac{X - \tilde{\theta}}{\sigma} \right) d\tilde{\theta} - Q(\underline{r}) \left[1 - \Psi \left(\frac{X - \theta^*}{\sigma} \right) + \Psi \left(\frac{X - \theta}{\sigma} \right) \right] = 0.$$

From (27), $X = \sigma \Psi^{-1}(B) + \theta$. Replacing this into the above, and integrating the first term by parts, we conclude that condition (5) can be restated as follows:

$$\begin{aligned} & (Z(\theta^*, \underline{r}) - Q(\underline{r})) \left(1 - \Psi \left(\Psi^{-1}(B) + \frac{\theta - \theta^*}{\sigma} \right) \right) \\ & + \int_{-\infty}^{\theta^*} \left(-\frac{dZ(\tilde{\theta}, \underline{r})}{d\tilde{\theta}} \right) \left(1 - \Psi \left(\Psi^{-1}(B) + \frac{\theta - \tilde{\theta}}{\sigma} \right) \right) d\tilde{\theta} - Q(\underline{r})B = 0. \end{aligned} \quad (28)$$

Note that the left-hand side of the above equation is decreasing in B . This guarantees that the above equation admits a unique solution for B —and therefore that (5) admits a unique solution, as stated in the main text. Next, noting that the left-hand side of the above equation is increasing in σ . It then follows from the implicit function theorem that B is also increasing in σ , as claimed above. Finally, one can also use the above formula to verify that, for any $\theta > \theta^*$, $\lim_{\sigma \rightarrow 0^+} B(\theta^*, \theta; \sigma) = 0$, while for $\theta = \theta^*$, $\lim_{\sigma \rightarrow 0^+} B(\theta^*, \theta; \sigma) = (Z(\theta^*, \underline{r}) - Q(\underline{r})) / Z(\theta^*, \underline{r})$, which completes the proof of Property 3.

Property 4. Let $\delta > 0$ be the threshold in property 1 and σ' be the threshold in property 2. Define

$$\hat{\Theta} \equiv \{(\theta^*, \theta) : \theta^* \in [\theta^\#(0^+) + \delta, \bar{\theta}], \theta \in [\theta^*, \bar{\theta}]\}.$$

There exists $0 < \sigma_1 < \sigma'$ such that, for any $\sigma < \sigma_1$, any $(\theta^*, \theta) \in \hat{\Theta}$,

$$G(\theta; \theta^*, \sigma) < 0 < R(\theta, \underline{r}, B(\theta^*, \theta; \sigma))$$

where, for any $\sigma > 0$, any (θ^*, θ) with $\theta \geq \theta^* \geq \underline{\theta}$,

$$G(\theta; \theta^*, \sigma) \equiv W(\theta, \hat{r}(\theta^*), 0) - W(\theta, \underline{r}, B(\theta^*, \theta; \sigma)).$$

Using property 3, for any $(\theta^*, \theta) \in \hat{\Theta}$,

$$\lim_{\sigma \rightarrow 0^+} R(\theta, \underline{r}, B(\theta^*, \theta; \sigma)) = \begin{cases} R(\theta, \underline{r}, 0) > 0 & \text{if } \theta > \theta^* \\ R(\theta^*, \underline{r}, (Z(\theta^*, \underline{r}) - Q(\underline{r})) / Z(\theta^*, \underline{r})) > 0 & \text{if } \theta = \theta^* \end{cases}$$

where the second inequality follows from the fact that the function $R(\theta^*, \underline{r}, (Z(\theta^*, \underline{r}) - Q(\underline{r})) / Z(\theta^*, \underline{r}))$ is strictly increasing in θ^* and equal to zero at $\theta^* = \theta^\#(0^+)$ (this follows from Lemma 1) along with the fact that $\theta^* > \theta^\#(0^+) + \delta$.

Likewise, for any $(\theta^*, \theta) \in \hat{\Theta}$,

$$\lim_{\sigma \rightarrow 0^+} G(\theta; \theta^*, \sigma) = \begin{cases} W(\theta, \hat{r}(\theta^*), 0) - W(\theta, \underline{r}, 0) < 0 & \text{if } \theta > \theta^* \\ L(\theta^*, \underline{r}) - W(\theta^*, \underline{r}, (Z(\theta^*, \underline{r}) - Q(\underline{r})) / Z(\theta^*, \underline{r})) < 0 & \text{if } \theta = \theta^* \end{cases}$$

where the second inequality follows from the fact that $R(\theta^*, \underline{r}, (Z(\theta^*, \underline{r}) - Q(\underline{r})) / Z(\theta^*, \underline{r})) > 0$ implies $W(\theta^*, \underline{r}, (Z(\theta^*, \underline{r}) - Q(\underline{r})) / Z(\theta^*, \underline{r})) > L(\theta^*, \underline{r})$.

Next, note that, for any $\sigma > 0$, the functions $R(\cdot, \underline{r}, B(\cdot, \cdot; \sigma))$ and $G(\cdot; \cdot, \sigma)$ are continuous on $\hat{\Theta}$ and that, for any $(\theta^*, \theta) \in \hat{\Theta}$, they are monotone in σ (by the monotonicity of B in σ). Because $\hat{\Theta}$ is compact, the result then follows from Dini's theorem.

Property 5. Let $\delta > 0$ be the threshold in property 1 and $\sigma_1 > 0$ the threshold in property 4. For any $\sigma < \sigma_1$,

$$G(\theta; \theta^*, \sigma) < 0 < R(\theta, \underline{r}, B(\theta^*, \theta; \sigma))$$

for any (θ^*, θ) such that either (a) $\theta^* \in [\theta^\#(0^+) + \delta, \bar{\theta}]$ and $\theta > \bar{\theta}$, or (b) $\theta^* > \bar{\theta}$ and $\theta \geq \theta^*$.

The result follows directly from properties 2 and 4.

Together, the above properties imply that for any $\sigma < \sigma_1$, $s > r^\#(\sigma) + \varepsilon$, there exists no pair (θ_s^*, θ_s'') such that θ_s^* satisfies conditions (8) and $\theta_s'' \geq \theta_s^*$ satisfies either condition (9) or condition (10). From Proposition (1) this means that, for any $\sigma < \sigma_1$, $\mathcal{E}(s; \sigma) = \emptyset$ for any $s > r^\#(\sigma) + \varepsilon$.

Part (ii). Let $\varepsilon > 0$ be the same as in part (i). By the continuity and strict monotonicity of the function $\hat{r}(\theta)$ over $[\underline{r}, +\infty)$, together with the continuity of the $\theta^\#(\sigma)$ function, there exist $\sigma^+ > 0$ and $\varepsilon^+ \in (0, \varepsilon)$ such that, for any $\sigma < \sigma^+$ $\theta_s^* > \theta^\#(\sigma) + \varepsilon$ implies that $s > r^\#(\sigma) + \varepsilon^+$. The result in part (i) then implies that there exists a $\sigma_2 \in (0, \sigma^+)$ such that, for any $\sigma < \sigma_2$, $\mathcal{E}(s; \sigma) \neq \emptyset$ only if $\theta_s^* \leq \theta^\#(\sigma) + \varepsilon$.

Part (iii). Take the same $\varepsilon > 0$ as in parts (i) and (ii) and let $r = \underline{r} + \varepsilon$. By the limit condition that $\lim_{\theta \rightarrow +\infty} [W(\theta, r, 0) - W(\theta, \underline{r}, 1)] < 0$, there exists a $\bar{\theta} > \bar{\theta}$ such that, for any $\theta > \bar{\theta}$, any A , any $s \geq \underline{r} + \varepsilon$, $R(\theta, \underline{r}, A) > 0$ and $W(\theta, s, 0) < W(\theta, \underline{r}, A)$. This means that, for any $s \geq \underline{r} + \varepsilon$, if a pair (θ_s^*, θ_s'') exists such that θ_s^* satisfies condition (8) and $\theta_s'' \geq \theta_s^*$ satisfies either condition (9) or condition (10), then necessarily $\theta_s^*, \theta_s'' \leq \bar{\theta}$. In turn, this also means that, for any $s \geq \underline{r} + \varepsilon$, $\mathcal{E}(s; \sigma) \neq \emptyset$ only if $s \leq \hat{r}(\bar{\theta})$. Thus assume $\hat{r}(\bar{\theta}) \geq \underline{r} + \varepsilon$. For any $s \in [\underline{r} + \varepsilon, \hat{r}(\bar{\theta})]$, then let θ_s^* be the threshold corresponding to s , as defined in (8) and take any arbitrary $\theta > \theta_s^*$. Because $\lim_{\sigma \rightarrow 0^+} B(\theta_s^*, \theta; \sigma) = 0$, and because $B(\theta_s^*, \cdot; \sigma)$ is decreasing in θ (this can be seen from (28)), we then have that, for any $s \in [\underline{r} + \varepsilon, \hat{r}(\bar{\theta})]$, either there exists no $\theta_s'' \geq \theta_s^*$ that satisfies condition (9) or condition (10) when σ is small enough, in which case $\mathcal{E}(s; \sigma) \neq \emptyset$. Else, if such a θ_s'' exists, then the

highest $\theta_s'' \geq \theta_s^*$ that satisfies either condition (9) or condition (10, which is $\theta^{**}(\sigma)$), must converge to θ_s^* as $\sigma \rightarrow 0$.

For any $s \in [\underline{r} + \varepsilon, \hat{r}(\bar{\theta})]$ any σ , then let $\Delta_s(\sigma) = \theta^{**}(\sigma) - \theta_s^*$ if a $\theta_s'' \geq \theta_s^*$ that satisfies either condition (9) or condition (10) exists and $\Delta_s(\sigma) = 0$ otherwise. Because for any $\theta^*, \theta, \theta \geq \theta^*$, $B(\theta^*; \cdot)$ is increasing in σ , we then have that, for any $s \in [\underline{r} + \varepsilon, \hat{r}(\bar{\theta})]$, $\Delta_s(\sigma)$ is increasing in σ with $\lim_{\sigma \rightarrow 0+} \Delta_s(\sigma) = 0$. Because, for any $\sigma > 0$, $\Delta_s(\sigma)$ is clearly continuous in s over $[\underline{r} + \varepsilon, \hat{r}(\bar{\theta})]$, from Dini's theorem we then have that there exists a $\sigma_3 > 0$ such that, for $\sigma < \sigma_3$, any $s \in [\underline{r} + \varepsilon, \hat{r}(\bar{\theta})]$, either $\mathcal{E}(s; \sigma) = \emptyset$, or $\theta_s^{**}(\sigma) \leq \theta_s^* + \varepsilon$, which establish the result.

Combining the proofs of parts (i)-(iii), the result in the proposition then follows by letting $\bar{\sigma} = \min\{\sigma_1, \sigma_2, \sigma_3\}$. *Q.E.D.*

Proof of Proposition 10. Part (i) follows from Lemma A3 below, whose proof is similar to that of Proposition 5 in Angeletos, Hellwig and Pavan (2006), adapted to the payoff structure considered here. Part (ii) follows from Lemma A4.

Lemma A3. *Assume ψ is log-concave. For any $\sigma > 0$, any $s \in (\underline{r}, r^\#(\sigma)]$, there exists a nonempty set of types $\Theta_s(\sigma) \subset [\theta_s^*, \theta_s^{**}(\sigma)]$, with $\inf \Theta_s(\sigma) = \theta_s^*$, a threshold x_s^* , and an equilibrium in which*

$$r(\theta) = \begin{cases} s & \text{if } \theta \in \Theta_s(\sigma) \\ \underline{r} & \text{otherwise} \end{cases} \quad \text{and} \quad D(\theta) = \begin{cases} 1 & \text{if } \theta < \theta_s^* \\ 0 & \text{if } \theta > \theta_s^* \end{cases}$$

The equilibrium is sustained by the following strategy for the agents: (i) for $r = \underline{r}$, $a(x, r) = 1$ if and only if $x < x_s^$; for any $r \in (\underline{r}, s)$, $a(x, r) = 1$ irrespective of x ; for any $r \geq s$, $a(x, r) = 0$ irrespective of x .*

Proof of Lemma A3. The proof is in three steps. Steps 1 and 2 construct the set $\Theta_s(\sigma)$ and establish properties that are useful for step 3. Step 3 shows that there exists a system of beliefs that support the proposed strategies as part of an equilibrium.

Step 1. Fix $s \in (\underline{r}, r^\#(\sigma)]$ and let $\theta_s^* = \inf\{\theta \geq \underline{\theta} : W(\theta, s, 0) \geq L(\theta, \underline{r})\}$. Next, let $S : \mathbb{R} \rightarrow 2^{\mathbb{R}}$ and $m : \mathbb{R}^2 \rightarrow [0, 1]$ denote, respectively, the correspondence and the function defined as follows:

$$\begin{aligned} S(x) &\equiv \left\{ \theta > \underline{\theta} : W(\theta, s, 0) \geq U\left(\theta, \underline{r}, \Psi\left(\frac{x-\theta}{\sigma}\right)\right) \right\}, \\ m(x, x') &\equiv \frac{\int_{-\infty}^{\theta_s^*} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{x'-\tilde{\theta}}{\sigma}\right) d\tilde{\theta}}{1 - \int_{S(x)} \frac{1}{\sigma} \psi\left(\frac{x'-\tilde{\theta}}{\sigma}\right) d\tilde{\theta}} - Q(\underline{r}) \end{aligned}$$

The set $S(x)$ represents the set of types who prefer raising the policy to $r = s$ and then facing no attack to leaving the policy at $\underline{r}$, facing an attack of size $\Psi\left(\frac{x-\theta}{\sigma}\right)$, and then obtaining a payoff

$$U\left(\theta, \underline{r}, \Psi\left(\frac{x-\theta}{\sigma}\right)\right) = \begin{cases} W\left(\theta, \underline{r}, \Psi\left(\frac{x-\theta}{\sigma}\right)\right) & \text{if } R\left(\theta, \underline{r}, \Psi\left(\frac{x-\theta}{\sigma}\right)\right) > 0 \\ L(\theta, \underline{r}) & \text{if } R\left(\theta, \underline{r}, \Psi\left(\frac{x-\theta}{\sigma}\right)\right) \leq 0 \end{cases}$$

By definition, any type below θ_s^* prefers leaving the policy at $r = \underline{r}$ to raising the policy to $r = s$, irrespective of the size of the attack. It follows that $\theta_s^* \leq \inf S(x)$ for any x . In turn, $m(x, x')$ is the expected payoff from attacking for an agent with signal x' when he observes $\underline{r}$ and believes that regime change will occur if and only if $\theta < \theta_s^*$ and that the policy is $r(\theta) = \underline{r}$ if and only if $\theta \notin S(x)$.

Step 2 below shows that, when ψ is log-concave, then for any x , the function $m(x, \cdot)$ is non-increasing in x' . It then uses this property to show that either (i) there exists an $x_s^* \in \mathbb{R}$ such that $m(x_s^*, x_s^*) = 0$, or (ii) $m(x, x) > 0$ for all x , in which case we let $x_s^* = +\infty$. In either case, the triple $(x_s^*, \theta_s^*, \Theta_s(\sigma))$ with $\Theta_s(\sigma) = S(x_s^*)$ identifies an equilibrium for the fictitious game in which the policy maker is restricted to set $r \in \{\underline{r}, s\}$ and the agents are restricted not to attack when $r = s$. Step 3 concludes the proof by showing that the same triple identifies an equilibrium for the unrestricted game.

Step 2. Note that $S(x)$ is continuous in x , with $S(x_1) \subseteq S(x_2)$ for any $x_1 \leq x_2$ (this follows from the fact that the policy maker's payoff from not raising the policy declines with the aggressiveness of the agents' behavior). Also note that $m(x, x')$ is continuous in (x, x') and nondecreasing in x (by the monotonicity of $S(x)$). Below we show that, when ψ is log-concave, $m(x, x')$ is also non-increasing in x' . To see this, note that, for any $\theta' \leq \theta_s^*$, the probability that an agent with signal x' assigns to $\tilde{\theta} < \theta'$ when observing $r = \underline{r}$ is $\mu(\theta'|x', \underline{r}) = (1 + 1/M(x'))^{-1}$, where

$$M(x') \equiv \frac{1 - \Psi\left(\frac{x' - \theta'}{\sigma}\right)}{\int_{\theta'}^{\infty} \left[1 - I_x(\tilde{\theta})\right] \frac{1}{\sigma} \psi\left(\frac{x' - \tilde{\theta}}{\sigma}\right) d\tilde{\theta}}$$

with $I_x(\tilde{\theta}) = 1$ when $\tilde{\theta} \in S(x)$ and $I_x(\tilde{\theta}) = 0$ otherwise. It follows that $\mu(\theta'|x', \underline{r})$ is decreasing in x' if $d \ln M(x') / dx' < 0$ or, equivalently, if

$$\frac{\int_{-\infty}^{\theta'} \frac{1}{\sigma^2} \psi'\left(\frac{x' - \tilde{\theta}}{\sigma}\right) d\tilde{\theta}}{\int_{-\infty}^{\theta'} \frac{1}{\sigma} \psi\left(\frac{x' - \tilde{\theta}}{\sigma}\right) d\tilde{\theta}} - \frac{\int_{\theta'}^{\infty} \left[1 - I_x(\tilde{\theta})\right] \frac{1}{\sigma^2} \psi'\left(\frac{x' - \tilde{\theta}}{\sigma}\right) d\tilde{\theta}}{\int_{\theta'}^{\infty} \left[1 - I_x(\tilde{\theta})\right] \frac{1}{\sigma} \psi\left(\frac{x' - \tilde{\theta}}{\sigma}\right) d\tilde{\theta}} < 0. \quad (29)$$

Using the fact that $I_x(\tilde{\theta}) = 0$ for all $\tilde{\theta} \leq \theta'$, (29) is equivalent to

$$\mathbb{E}_{\tilde{\theta}} \left[\frac{\psi'\left(\frac{x' - \tilde{\theta}}{\sigma}\right)}{\psi\left(\frac{x' - \tilde{\theta}}{\sigma}\right)} \middle| \tilde{\theta} \leq \theta', x', \underline{r} \right] - \mathbb{E}_{\tilde{\theta}} \left[\frac{\psi'\left(\frac{x' - \tilde{\theta}}{\sigma}\right)}{\psi\left(\frac{x' - \tilde{\theta}}{\sigma}\right)} \middle| \tilde{\theta} > \theta', x', \underline{r} \right] < 0,$$

which holds true when ψ'/ψ is decreasing, i.e. when ψ is log-concave. That $Z(\cdot, \underline{r})$ is non-increasing together with the fact that $\mu(\theta'|x', \underline{r})$ is decreasing in x' for any $\theta' \leq \theta_s^*$ then implies that the expected payoff from attacking

$$m(x, x') = \frac{\int_{-\infty}^{\theta_s^*} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{x' - \tilde{\theta}}{\sigma}\right) d\tilde{\theta}}{1 - \Psi\left(\frac{x' - \theta_s^*}{\sigma}\right) + \int_{\theta_s^*}^{\infty} \left[1 - I_x(\tilde{\theta})\right] \frac{1}{\sigma} \psi\left(\frac{x' - \tilde{\theta}}{\sigma}\right) d\tilde{\theta}} - \underline{r}$$

is non-increasing in x' .⁴⁵ For future reference, also note that $\lim_{x' \rightarrow +\infty} m(x, x') < 0$.

Having established that the expected payoff from attacking $m(x, x')$ is continuous in (x, x') , nondecreasing in x and non-increasing in x' , next note that

$$S(x) \subseteq \bar{S} \equiv \{\theta : W(\theta, s, 0) \geq U(\theta, r, 1)\}$$

which in turn implies that, for any $(x, x') \in \mathbb{R}^2$,

$$\frac{\int_{-\infty}^{\theta_s^*} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{x' - \tilde{\theta}}{\sigma}\right) d\tilde{\theta}}{1 - \int_{\bar{S}} \frac{1}{\sigma} \psi\left(\frac{x' - \tilde{\theta}}{\sigma}\right) d\tilde{\theta}} - Q(\underline{r}) \geq m(x, x') \geq \int_{-\infty}^{\theta_s^*} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{x' - \tilde{\theta}}{\sigma}\right) d\tilde{\theta} - Q(\underline{r}).$$

It follows that, for any x , $m(x, x') \geq 0$ for all $x' \leq \hat{x}$, where $\hat{x} \in \mathbb{R}$ is the unique solution to

$$\int_{-\infty}^{\theta_s^*} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{\hat{x} - \tilde{\theta}}{\sigma}\right) d\tilde{\theta} = Q(\underline{r})$$

Now define the sequence $\{x_k\}_{k=0}^{\infty}$, with $x_k \in \mathbb{R} \cup \{+\infty\}$, as follows. For $k = 0$, let $x_0 \equiv \hat{x}$. For $k \geq 1$, let x_k be the solution to $m(x_{k-1}, x_k) = 0$ if $x_{k-1} < +\infty$; if, instead, $x_{k-1} = +\infty$, let $x_k \equiv \inf\{x' : m(x_{k-1}, x') \leq 0\}$ if $\{x' : m(x_{k-1}, x') \leq 0\} \neq \emptyset$ and $x_k \equiv +\infty$ otherwise. The fact that

$$m(\hat{x}, \hat{x}) \geq \int_{-\infty}^{\theta_s^*} Z(\tilde{\theta}, \underline{r}) \frac{1}{\sigma} \psi\left(\frac{\hat{x} - \tilde{\theta}}{\sigma}\right) d\tilde{\theta} - Q(\underline{r}) = 0$$

together with the continuity and monotonicities of m , ensures that this sequence is well defined and nondecreasing. It follows that either $\lim_{k \rightarrow \infty} x_k \in [\hat{x}, +\infty)$, or $\lim_{k \rightarrow \infty} x_k = +\infty$. In the former case, let $x_s^* = \lim_{k \rightarrow \infty} x_k$ and $\Theta_s(\sigma) = S(x_s^*)$; in the latter, let $x_s^* = +\infty$ and $\Theta_s(\sigma) = S(\infty) \equiv \bar{S}$.

Finally, to see that $\inf \Theta_s(\sigma) = \theta_s^*$ note that the threshold $\hat{x}$ defined above coincides with $X(\theta_s^*, \theta_s^*, \sigma)$. Because $\theta_s^* \leq \theta^\#(\sigma)$, from the results in the proof of Lemma A1 above, we then have that $R(\theta_s^*, \underline{r}, \Psi((X(\theta_s^*, \theta_s^*, \sigma) - \theta_s^*)/\sigma)) \leq 0$. Because $x_s^* \geq \hat{x}$, we then have that $R(\theta_s^*, \underline{r}, \Psi((x_s^* - \theta_s^*)/\sigma)) \leq 0$. Because $R(\theta, \underline{r}, A)$ is increasing in θ and decreasing in A and because $A(\theta, \underline{r}) = \Psi(\frac{x_s^* - \theta}{\sigma})$ is decreasing in θ , this in turn implies that there exists $\hat{\theta}_s \in [\theta_s^*, \theta_s^{**}(\sigma)]$ such that $R(\hat{\theta}_s, \underline{r}, \Psi((x_s^* - \hat{\theta}_s)/\sigma)) \leq 0$ if and only if $\theta \leq \hat{\theta}_s$. It follows that necessarily $r(\theta) = s$ for all $\theta \in (\theta_s^*, \hat{\theta}_s]$, thus establishing that $\inf \Theta_s(\sigma) = \theta_s^*$.

We conclude that the triple $(x_s^*, \theta_s^*, \Theta_s(\sigma))$ identifies an equilibrium for the fictitious game in which the policy maker is restricted to set $r \in \{\underline{r}, s\}$ and the agents are restricted not to attack when $r = s$.

Step 3. We now show how the triple $(x_s^*, \theta_s^*, \Theta_s(\sigma))$ of Step 2 also identifies an equilibrium for the unrestricted game. The proof here parallels that of Lemma A2. Below we simply show existence of beliefs for the agents satisfying conditions (19), (20) and (21).

⁴⁵To see this, note that, given any two absolutely continuous c.d.f.s F_1 and F_2 with $F_1(\theta) \leq F_2(\theta)$ for all $\theta < \theta_s^*$, and any non-increasing differentiable positive function $h : \mathbb{R} \rightarrow \mathbb{R}_+$, $\int_{-\infty}^{\theta_s^*} h(\tilde{\theta}) dF_1(\tilde{\theta}) \leq \int_{-\infty}^{\theta_s^*} h(\tilde{\theta}) dF_2(\tilde{\theta})$. Interpreting h as the payoff Z from regime change and F_1 and F_2 as the agent's posterior for two different signals x_1 and x_2 , with $x_1 \geq x_2$, gives the result.

When $r = \underline{r}$, beliefs are pinned down by Bayes rule and such that, for any $\theta \leq \theta_s^*$,

$$\mu(\theta|x, \underline{r}) = \frac{1 - \Psi\left(\frac{x-\theta}{\sigma}\right)}{1 - \int_{\Theta_s(\sigma)} \frac{1}{\sigma} \psi\left(\frac{x-\tilde{\theta}}{\sigma}\right) d\tilde{\theta}},$$

while for any $\theta \in (\theta_s^*, \hat{\theta}_s)$,⁴⁶

$$\mu(\theta|x, \underline{r}) = \frac{1 - \Psi\left(\frac{x-\theta_s^*}{\sigma}\right)}{1 - \int_{\Theta_s(\sigma)} \frac{1}{\sigma} \psi\left(\frac{x-\tilde{\theta}}{\sigma}\right) d\tilde{\theta}}.$$

As shown above, these beliefs are decreasing in x . By the definition of x_s^* , it then follows that condition (19) is satisfied when $r = \underline{r}$. Next consider $r = s$. Again, in this case beliefs are pinned down by Bayes rule and such that $\mu(\underline{\theta}|x, s) = 0$, all x , in which case condition (21) is clearly satisfied. Finally, whenever $r \notin \{\underline{r}, s\}$, there exist an arbitrarily large set of out-of-equilibrium beliefs that satisfy (20) and (21) — see the construction in Lemmas A0 and A2 above. ■

Lemma A4. *Suppose SCC holds and ψ is log-concave. Then for any $\sigma > 0$, any $s \in (\underline{r}, r^\#(\sigma)]$, any equilibrium in $\mathcal{E}(s; \sigma)$ is such that $r(\theta) = s$ for all $\theta \in (\theta_s^*, \theta_s^{**}(\sigma))$.*

Proof of Lemma A4. The result for $s = r^\#(\sigma)$ follows directly from the fact that, when SCC holds, then $\theta_s^* = \theta^\#(\sigma) = \theta_s^{**}(\sigma)$. That $\theta_s^* = \theta^\#(\sigma)$ is immediate. To see that $\theta_s^{**}(\sigma) = \theta^\#(\sigma)$, recall that, by the properties of the pooling equilibria, for any $\theta > \theta^\#(\sigma)$,

$$R\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta^\#(\sigma), \theta^\#(\sigma); \sigma) - \theta}{\sigma}\right)\right) > 0$$

which implies that

$$U\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta^\#(\sigma), \theta^\#(\sigma); \sigma) - \theta}{\sigma}\right)\right) = W\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta^\#(\sigma), \theta^\#(\sigma); \sigma) - \theta}{\sigma}\right)\right) > L(\theta, \underline{r}). \quad (30)$$

By continuity, when applied to $\theta = \theta^\#(\sigma)$, (30) implies that

$$W\left(\theta^\#(\sigma), r^\#(\sigma), 0\right) = L(\theta^\#(\sigma), \underline{r}) \leq W\left(\theta^\#(\sigma), \underline{r}, \Psi\left(\frac{X(\theta^\#(\sigma), \theta^\#(\sigma); \sigma) - \theta^\#(\sigma)}{\sigma}\right)\right)$$

Under SCC, this means that, for any $\theta > \theta^\#(\sigma)$,

$$\begin{aligned} U\left(\theta, r^\#(\sigma), 0\right) &= W\left(\theta, r^\#(\sigma), 0\right) \leq U\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta^\#(\sigma), \theta^\#(\sigma); \sigma) - \theta^\#(\sigma)}{\sigma}\right)\right) \\ &= W\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta^\#(\sigma), \theta^\#(\sigma); \sigma) - \theta^\#(\sigma)}{\sigma}\right)\right) \end{aligned}$$

Because $\Psi\left(\frac{X(\theta^\#(\sigma), \theta^\#(\sigma); \sigma) - \theta^\#(\sigma)}{\sigma}\right) > \Psi\left(\frac{X(\theta^\#(\sigma), \theta; \sigma) - \theta}{\sigma}\right)$, as shown in the proof of Proposition 2, this means that, for any $\theta > \theta^\#(\sigma)$, $R\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta^\#(\sigma), \theta; \sigma) - \theta}{\sigma}\right)\right) > 0$ and $W(\theta, r^\#(\sigma), 0) \leq$

⁴⁶Recall that all types $\theta \in (\theta_s^*, \hat{\theta}_s)$ necessarily raise the policy to $r = s$.

$W\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta^\#(\sigma), \theta; \sigma) - \theta}{\sigma}\right)\right)$, which implies that the only $\theta_s'' \geq \theta_s^* = \theta^\#(\sigma)$ that possibly satisfies (9) or (10) is $\theta_s'' = \theta^\#(\sigma)$, which means that $\theta_s^{**}(\sigma) = \theta_s^* = \theta^\#(\sigma)$ for $s = r^\#(\sigma)$.

Thus consider $s \in (\underline{r}, r^\#(\sigma))$. From the proof of Lemma 4, $a(x, \underline{r}) = 0$ for all $x > X(\theta_s^*, \theta_s^{**}(\sigma); \sigma)$, while from the proof of Lemma 5, $a(x, \underline{r}) = 1$ for all $x < X(\theta_s^*, \theta_s^*; \sigma)$. It follows that $\Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma)) - \theta}{\sigma}\right) \geq A(\theta, \underline{r}) \geq \Psi\left(\frac{X(\theta_s^*, \theta_s^*) - \theta}{\sigma}\right)$ for all θ . By the fact that $R\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta, \theta; \sigma) - \theta}{\sigma}\right)\right) < 0$ for $\theta < \theta^\#(\sigma)$, we then have that $R\left(\theta_s^*, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma)) - \theta_s^*}{\sigma}\right)\right) < 0$, while by the fact that $\theta_s^{**}(\sigma)$ satisfies condition (9) or (10), we have that $R\left(\theta_s^{**}(s), \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta_s^{**}(\sigma)) - \theta_s^{**}(\sigma)}{\sigma}\right)\right) \geq 0$. Combining, we have that

$$R(\theta_s^*, \underline{r}, A(\theta_s^*, \underline{r})) < 0 \leq R(\theta_s^{**}(s), \underline{r}, A(\theta_s^{**}(s), \underline{r})) \quad (31)$$

Now the fact that the noise distribution ψ is log-concave implies that, after observing $r = \underline{r}$, irrespective of the shape of the equilibrium policy $r(\theta)$ in the region $[\theta_s^*, \theta_s^{**}(\sigma)]$ of possible intervention, the aggregate attack $A(\theta, \underline{r})$ is monotone in θ — see the proof of Lemma A3. Condition (31), together with the monotonicity of $A(\theta, \underline{r})$ in θ and the property that $R_\theta > 0 > R_A$ then ensures that there exists a unique $\hat{\theta}_s \in [\theta_s^*, \theta_s^{**}(\sigma)]$ such that $R(\theta, \underline{r}, A(\theta, \underline{r})) \leq 0$ if and only if $\theta \leq \hat{\theta}_s$. Now if $\hat{\theta}_s = \theta_s^{**}(\sigma)$, then obviously $r(\theta) = s$ for all $\theta \in (\theta_s^*, \theta_s^{**}(\sigma)]$. Thus suppose that $\hat{\theta}_s < \theta_s^{**}(\sigma)$, which means $\theta_s^{**}(\sigma)$ satisfies condition (9). Now let $\theta_s^{\text{sup}} \equiv \sup\{\theta : r(\theta) = s\}$. Clearly, $\theta_s^{\text{sup}} \geq \hat{\theta}_s$; if $\theta_s^{\text{sup}} < \hat{\theta}_s$, then types $\theta \in (\theta_s^{\text{sup}}, \hat{\theta}_s]$ would be better off by raising the policy to $r = s$ rather than leaving the policy at $\underline{r}$ and then facing regime change.

We now show that all types $\theta \in (\theta_s^*, \theta_s^{\text{sup}}]$ necessarily raise the policy to $r = s$. This is immediate when $\theta_s^{\text{sup}} = \hat{\theta}_s$. Thus suppose $\theta_s^{\text{sup}} > \hat{\theta}_s$. Because, in this case, $R(\theta_s^{\text{sup}}, \underline{r}, A(\theta_s^{\text{sup}}, \underline{r})) > 0$, then θ_s^{sup} must satisfy

$$W(\theta_s^{\text{sup}}, s, 0) = W(\theta_s^{\text{sup}}, \underline{r}, A(\theta_s^{\text{sup}}, \underline{r}))$$

Now let

$$h(\theta) \equiv W(\theta, s, 0) - W(\theta, \underline{r}, A(\theta, \underline{r}))$$

Note that SCC, along with the monotonicity of $A(\cdot, \underline{r})$, implies that $h(\theta) > 0$ for all $\theta \in (\theta_s^*, \theta_s^{\text{sup}})$, which in turn implies that all $\theta \in (\theta_s^*, \theta_s^{\text{sup}})$ necessarily raise the policy to $r = s$. But then necessarily

$$A(\theta, \underline{r}) = \Psi\left(\frac{X(\theta_s^*, \theta_s^{\text{sup}}; \sigma) - \theta}{\sigma}\right)$$

This means that θ_s^{sup} must satisfy condition (9). Now recall that $\theta_s^{**}(\sigma)$ is the highest θ_s'' that satisfies (9). That $\theta_s^{\text{sup}} = \theta_s^{**}(\sigma)$ in turn follows from the fact that, under SCC, the function

$$G(\theta; \theta_s^*, \sigma) \equiv W(\theta, s, 0) - W\left(\theta, \underline{r}, \Psi\left(\frac{X(\theta_s^*, \theta; \sigma) - \theta}{\sigma}\right)\right)$$

is strictly negative for all $\theta \geq \theta_s^{\text{sup}}$ which implies that $\theta_s^{\text{sup}} = \theta_s^{**}(\sigma)$. ■ *Q.E.D.*

References

- [1] Angeletos, George-Marios, Christian Hellwig, and Alessandro Pavan (2006) "Signaling in a Global Game: Coordination and Policy Traps," *Journal of Political Economy* 114(3), 452-484.
- [2] Angeletos, George-Marios, Christian Hellwig, and Alessandro Pavan (2007) "Dynamic Global Games of Regime Change: Learning, Multiplicity, and the Timing of Attacks," *Econometrica*, 75(3), 711-756.
- [3] Angeletos, George-Marios, and Iván Werning (2006), "Crises and Prices: Information Aggregation, Multiplicity and Volatility," *American Economic Review* 96, 1721-1737.
- [4] Atkeson, Andrew (2001): "Discussion on Morris and Shin," *NBER Macroeconomics Annual* 2000.
- [5] Calvo, Guillermo (1988) "Servicing the Public Debt: the Role of Expectations." *American Economic Review* 78, 647-661.
- [6] Carlsson, Hans, and Eric van Damme (1993), "Global Games and Equilibrium Selection," *Econometrica* 61, 989-1018.
- [7] Ciliberto, F. and E. Tamer, 2009, "Market Structure and Multiple Equilibria in Airline Markets," *Econometrica*, 77, 1791-1828.
- [8] Chamley, Christophe (1999), "Coordinating Regime Switches," *Quarterly Journal of Economics* 114, 869-905.
- [9] Chamley, Christophe (2003), "Dynamic Speculative Attacks," *American Economic Review*, 93(3), 603-621.
- [10] Chari V.V. & Patrick J. Kehoe (2003) "Hot Money," *Journal of Political Economy*, 111(6), 1262-1292.
- [11] Chassang, Sylvain (2010), "Fear of Mis-coordination and the Robustness of Cooperation in Dynamic Global Games with Exit, *Econometrica*.
- [12] Cho, In-Koo, and David Kreps (1987) "Signaling Games and Stable Equilibria," *Quarterly Journal of Economics* 102, 179-222.
- [13] Corsetti, Giancarlo, Amil Dasgupta, Stephen Morris and Hyun Song Shin (2004) "Does One Soros Make a Difference? A Theory of Currency Crises with Large and Small Traders" *The Review of Economic Studies*, 71(1), 87-114.

- [14] Corsetti, Giancarlo, Bernardo Guimaraes, and Nouriel Roubini (2006), "International Lending of Last Resort and Moral Hazard: A Model of IMF's Catalytic Finance," *Journal of Monetary Economics* 53, 441-471.
- [15] Dasgupta, Amil (2006): "Coordination and Delay in Global Games" *Journal of Economic Theory*, forthcoming.
- [16] Edmond, Chris (2011): "Information Manipulation, Coordination and Regime Change," working paper, NYU Stern School of Business.
- [17] Drazen, Allan (2000), "Interest-rate and Borrowing Defense Against Speculative Attack," *Carnegie-Rochester Conference Series on Public Policy* 53, 303-248.
- [18] Drazen, Allan, and Stefan Hubrich (2005), "A Simple Test of the Effect of Exchange Rate Defense," working paper, University of Maryland.
- [19] Frankel, David, Stephen Morris, and Ady Pauzner (2003), "Equilibrium Selection in Global Games with Strategic Complementarities," *Journal of Economic Theory*, 108, 1-44.
- [20] Goldstein, Itay, and Ady Pauzner (2005), "Demand Deposit Contracts and the Probability of Bank Runs," *Journal of Finance* 60, 1293-1328.
- [21] Goldstein, Itay, Emre Ozdenoren, and Kathy Yuan, (2011), "Learning and Complementarities in Speculative Attacks", the *Review of Economic Studies*, vol. 78(1), 263-292.
- [22] Grieco, Paul (2010), "Discrete Games with Flexible Information Structures: An Application to Local Grocery Markets," working paper, Pennsylvania State University.
- [23] Guimaraes, Bernardo, and Stephen Morris (2006), "Risk and Wealth in a Model of Self-fulfilling Currency Attacks," working paper, LSE and Yale University.
- [24] Hellwig, Christian, Arijit Mukherji, and Aleh Tsyvinski (2006), "Self-Fulfilling Currency Crises: the Role of Interest Rates," *American Economic Review*, 96, 1770-1788.
- [25] Jeanne, Olivier, and Charles Wyplosz (2001), "The International Lender of Last Resort: How Large is Large Enough?" working paper #8381, NBER.
- [26] Kraay, Aart (2001), "Do High Interest Rates Defend Currencies During Speculative Attack?," World Bank mimeo.
- [27] Dewan, Torun, and David Myatt (2007), "Leading the Party: Coordination, Direction, and Communication," *American Political Science Review* 101(4), 827-845.

- [28] Morris, Stephen, and Hyun Song Shin (1998), “Unique Equilibrium in a Model of Self-Fulfilling Currency Attacks,” *American Economic Review* 88, 587-597.
- [29] Morris, Stephen, and Hyun Song Shin (2001), “Rethinking Multiple Equilibria in Macroeconomics,” *NBER Macroeconomics Annual 2000*.
- [30] Morris, Stephen, and Hyun Song Shin (2003), “Global Games—Theory and Applications,” in *Advances in Economics and Econometrics*, 8th World Congress of the Econometric Society (M. Dewatripont, L. Hansen, and S. Turnovsky, eds.), Cambridge University Press.
- [31] Morris, Stephen, and Hyun Song Shin (2004), “Liquidity Black Holes,” *Review of Finance* 8, 1-18.
- [32] Morris, Stephen, and Hyun Song Shin (2006a), “Endogenous Public Signals and Coordination,” working paper, Princeton University.
- [33] Morris, Stephen, and Hyun Song Shin (2006b), “Catalytic Finance,” *Journal of International Economics* 70, 161-177.
- [34] Obstfeld, Maurice (1996), “Models of Currency Crises with Self-Fulfilling Features,” *European Economic Review* 40, 1037-47.
- [35] Ozdenoren, Emre, and Kathy Yuan (2008), “Feedback Effects and Asset Prices,” *Journal of Finance* 63(4), 1939-1975.
- [36] Penta, Antonio (2012), “Higher Order Uncertainty and Information: Static and Dynamic Games,” *Econometrica* 80, 631-660.
- [37] Rochet, Jean-Charles, and Xavier Vives (2004), “Coordination Failures and the Lender of Last Resort: Was Bagehot Right After All?” *Journal of the European Economic Association* 2-6, 1116-1147.
- [38] Tarashev, Nikola, (2007), “Currency Crises and the Informational Role of Interest Rates,” *Journal of the European Economic Association*, 5(4), 1-36.
- [39] Tamer, Elie (2003), “Incomplete Simultaneous Discrete Response Model with Multiple Equilibria,” *Review of Economic Studies*, 147–165.
- [40] Weinstein, Jonathan, and Muhamet Yildiz (2007), “A Structure Theorem for Rationalizability with Application to Robust Predictions of Refinements”, *Econometrica*, 75(2), 365-400.
- [41] Zettelmeyer, Jeromin (2000), “Can Official Crisis Lending be Counterproductive in the Short Run?” *Economic Notes* 29, 12-29.

- [42] Zwart, Sanne (2007) “The mixed blessing of IMF intervention: Signalling versus liquidity support,” *Journal of Financial Intermediation*, 3, 149–174.