

Coordination with sequential information acquisition

XIAOYE LIAO

Department of Economics and Data Science, New Uzbekistan University

MICHAL SZKUP

Department of Economics, University of British Columbia

We investigate how differences in initial beliefs and sequential information choices affect the likelihood of coordination failure. To do so, we embed interim information acquisition (i.e., information acquisition after observing initial private information) into a standard global game mode with a normal information structure and improper prior. We find that the likelihood of coordination on welfare-inferior equilibrium is invariant to precision, cost, and availability of information. We show that agents' information choices feature two-sided inefficiency where too many agents with high posteriors and too few agents with low posteriors acquire information. Instead, under efficient information choices, the likelihood of coordination failure vanishes as the number of signals that agents can acquire tends to infinity. Unfortunately, efficient information choices are not implementable unless policymakers can observe agents' private information.

KEYWORDS. Coordination failure, sequential information acquisition, global games, efficiency.

JEL CLASSIFICATION. C72, D83, D84.

1. INTRODUCTION

Many undesirable outcomes in economics are thought to be a result of coordination failure, that is, the inability of agents to coordinate on socially preferable actions. Examples include bank runs, credit freezes, currency crises, rollover crises, and sovereign debt crises, among many others.¹ As stressed initially by [Morris and Shin \(1998, 2002\)](#), agents' ability to coordinate on a given action depends crucially on the information agents possess, and, thus, the information structure plays a key role in determining coordination outcomes. A pertinent question then is how agents' information acquisition choices, which endogenously determine information structure, affect coordination outcomes.

While there exists a large literature on information acquisition in coordination games (see, e.g., [Hellwig and Veldkamp \(2009\)](#), [Myatt and Wallace \(2012\)](#), [Colombo,](#)

Xiaoye Liao: xl775@nyu.edu

Michał Szkup: michal.szakup@ubc.ca

We thank the two anonymous referees for comments and suggestions that greatly improved the paper. Michał Szkup gratefully acknowledges the financial support of the Social Sciences and Humanities Research Council of Canada (SSHRC).

¹See [Cooper and John \(1988\)](#), [Cooper \(1999\)](#), or [Angeletos and Lian \(2016\)](#) for more examples of coordination failures that arise in the context of macroeconomics.

Femminis, and Pavan (2014), Szkup and Trevino (2015), or Yang (2015)), this literature assumes that all agents share the same initial beliefs and, hence, the same motives to acquire information. However, decision-makers likely exhibit differences in their incentives to acquire information that arise from differences in their initial beliefs. These differences can arise due to having different prior experiences, having access to different prior sources of information, or having different interpretations of common information. Furthermore, even if decision-makers initially share the same incentives to acquire information, the sequential way in which information is typically acquired and processed implies that agents' interim beliefs, and hence interim information acquisition incentives, will differ. Motivated by these observations, in this paper, we ask how differences in initial information and sequential information acquisition affect coordination outcomes.

To answer this question, we introduce heterogeneity in incentives to acquire information and sequential information acquisition in a global game model of regime change. Global games of regime change are a natural setting for studying how heterogeneity in incentives to acquire information and sequential information acquisition affect coordination outcomes. First, they have been used to analyze a plethora of economic phenomena featuring coordination failure (Morris and Shin (2003) and Angeletos and Lian (2016) provide overviews of these applications). Second, information structure plays a key role in determining outcomes in global games. Finally, the incentive to sequentially acquire information naturally arises in these setups.

We consider a binary-action global game model with a Gaussian information structure where a continuum of agents decides whether to attack a regime as in Vives (2014). The regime changes if it is weak enough and/or if sufficiently many agents attack it. Agents do not observe the strength of the regime, but each of them observes a free private signal about it. In addition, before deciding whether to attack the regime but after observing the initial signal, agents can acquire an additional signal at a cost. The presence of the initial private signal introduces heterogeneity in agents' incentives to acquire information. Note that agents make their information choices after observing the initial signal, which can be interpreted as a simple form of sequential information.²

We use the benchmark model to investigate how differences in agents' initial information affect the coordination outcome. In particular, we first characterize agents' optimal information choices. We then investigate how their information decisions affect the equilibrium coordination outcome and the incidence of coordination failure. We also analyze how a reduction in the cost of information or changes in information precision affect coordination. Furthermore, we investigate the (in)efficiency of equilibrium information choices. Finally, we investigate how the answers to these questions change when agents can sequentially acquire up to $N - 1$ additional signals.

The main finding of our equilibrium analysis is an *invariance result*, which states that under relatively general conditions, the equilibrium regime-change threshold, and

²While the benchmark model features only two signals, we also consider an extended model with fully-fledged sequential information acquisition, where agents can sequentially acquire up to $N - 1$ additional signals, $N > 2$.

hence the extent of coordination failure, is unaffected by the parameters governing information choices. In other words, changes in the cost, precision, and number of available signals (in the extended model) leave the equilibrium regime-change threshold unchanged. Thus, the invariance result suggests that the likelihood of bank runs, sovereign debt crises, currency crises, and other outcomes thought to be a result of coordination failure is unaffected by the decrease in the cost of information or an increase in its precision and availability. Nevertheless, it is worth stressing that our results do not imply that information choices do not matter. Indeed, changes in the parameters governing sequential information choices affect both agents' individual decisions and the proportion of agents attacking the regime (away from the regime-change threshold), with higher information availability leading to a higher proportion of agents making "correct" decisions in those states.³ In other words, it is only the extent of coordination failure that is invariant and not the equilibrium itself.

We then investigate the efficiency of agents' information decisions. In particular, we consider a planner who would like to minimize the extent of coordination failure and who can control agents' information choices but not their decisions whether to attack the regime. We find that the planner can induce a substantially lower regime-change threshold by controlling agents' information choices. The sequential information acquisition policy prescribed by the planner is simple and intuitive: it stipulates that agents acquire an additional signal only if, based on their initial signals, they would not attack the regime. Comparing the planner's information acquisition strategy with the one used by agents in equilibrium, we see that the equilibrium information choices exhibit two-sided inefficiency with too many agents with high posterior beliefs and too few agents with low posterior beliefs acquiring further information. Furthermore, in the extended model with N signals, we show that as the number of signals increases to infinity, coordination failure vanishes under the planner's strategy. Unfortunately, we also find that the planner's policy is not implementable unless policymakers can observe agents' beliefs (or, equivalently, agents' private signals).

In the final part of the paper, we discuss several extensions of the baseline model. We show that the invariance result continues to hold when the precision and cost of additional signals vary with the number of already acquired signals and when agents are ex ante heterogeneous with respect to their payoffs, cost, precision of information, and the number of available signals. We also extend our analysis to the case where, at each point, agents can choose an information source from which to acquire information (where information sources differ in terms of the precision of information they provide and its cost), and to the case where agents can acquire information about other agents' past decisions whether to attack the regime (as in [Dasgupta \(2007\)](#)). However, we emphasize that the invariance result does depend on one particular assumption, namely the use of improper prior. Nevertheless, the invariance result is a good approximation for the case where agents start with a proper but relatively diffuse prior. Moreover, the improper prior has been popular in global games literature as it attempts to capture the unpredictability of outcomes, a common feature of tumultuous periods that precede regime

³Namely, attacking the regime when it will change and vice versa.

changes. Thus, the invariance result is an important benchmark for understanding the impact of changes in the cost, precision, and availability of information.

Related literature

Our paper contributes to the large and growing literature on the role of information in coordination games. This literature was initiated by the seminal contributions of [Morris and Shin \(1998, 2002\)](#), who analyzed the role of information in the context of global games and quadratic-Gaussian setups, respectively. Since then, their initial insights have been further extended and generalized in many directions in both settings under the assumption of exogenous information structure. See, for example, [Angeletos and Pavan \(2007\)](#) and [Ui and Yoshizawa \(2015\)](#) for further analysis of quadratic-Gaussian setups and [Hellwig \(2002\)](#), [Morris and Shin \(2004\)](#), [Guimaraes and Morris \(2007\)](#), and [Iachan and Nenov \(2015\)](#) for contributions to the global games literature.

We contribute to the more recent literature that considers information acquisition. The closest to our work are [Hellwig and Veldkamp \(2009\)](#), [Myatt and Wallace \(2012\)](#), and [Colombo, Femminis, and Pavan \(2014\)](#), who study ex ante information acquisition (i.e., one-time information choices based on common prior beliefs) in quadratic-Gaussian setups and [Szkup and Trevino \(2015\)](#), [Yang \(2015\)](#), and [Ahnert and Kakhbod \(2017\)](#), who analyze ex ante information acquisition in global games of regime change. In contrast, we focus on interim information acquisition, where agents decide whether to acquire additional information after observing initial private signals. Therefore, in our model, agents have heterogeneous incentives to acquire information and, if $N > 2$, choose how much information to acquire sequentially. As we show, these differences have important consequences for the conclusions we reach.

Our work is also related to papers that analyze the dynamic arrival of information in the context of global games (see [Angeletos, Hellwig, and Pavan \(2007\)](#), [Steiner \(2008\)](#), [Dasgupta, Steiner, and Stewart \(2012\)](#), and [Mathevet and Steiner \(2013\)](#)). However, in those papers, the focus is on the dynamics of agents' action choices, and private information arrival is exogenous and independent of agents' choices. Also related is [Dasgupta \(2007\)](#), who analyzes how the option to delay decisions affects equilibrium in a two-period global game model, where, by delaying their actions, agents can observe an accurate additional signal but at the cost of a lower future payoff from successful risky action. In contrast, in our model, agents face a direct cost of acquiring further private signals (and, in the extended model, can sequentially acquire multiple signals). As a consequence, the optimal information acquisition choices and, thus, equilibrium characterization differ in these two settings. Thus, more broadly, our paper highlights the differences between delay and interim information acquisition.

Individual decision problems with sequential information choices have a long tradition in statistics (see [DeGroot \(2005\)](#) for a summary of this early literature). Variants of these models have been analyzed in various economics contexts, typically under a Bayesian framework, by [Roberts and Weitzman \(1981\)](#), [Moscarini and Smith \(2001\)](#), and [Ke and Villas-Boas \(2019\)](#). Many key ideas of sequential analysis were applied to study bandit problems (see, e.g., [Rothschild \(1974\)](#) and [Gittins \(1979\)](#)), search and learning problems (see, e.g., [McCall \(1970\)](#) and [Weitzman \(1979\)](#)), and strategic information

transmission (see, e.g., [Liao \(2021\)](#)). In contrast, we consider sequential learning in a strategic environment. We also differ from those papers by considering a finite-horizon discrete-time problem.

2. MODEL

We consider a general binary-action model of regime change with incomplete information, extended to feature sequential information choices, in which agents decide whether to attack or support the status quo. While we cast our model in neutral terms referring to agents' final actions as supporting or attacking the regime, the model admits various interpretations including, among others, a model of bank runs, currency crises, and sovereign debt crises.⁴

2.1 Setup

The economy is initially in a status quo regime and is characterized by an unobservable state $\theta \in \mathbf{R}$, referred to as the *fundamental*, which captures the inverse of the current regime's strength (with a higher θ corresponding to a weaker regime). There is a continuum of risk-neutral agents indexed by $i \in [0, 1]$. Each agent i makes a binary decision a_i whether to attack ($a_i = 1$) or to support ($a_i = 0$) the regime, referred to as the *final decision*. Agents' aggregate attack is denoted by p (i.e., $p = \int_0^1 a_i \, di$).

The attack either succeeds, in which case the regime changes, or it fails, in which case the regime survives. The attack succeeds if $R(\theta, p) \geq 0$, where $R(\cdot, \cdot)$ is continuously differentiable and strictly increasing in both arguments, so that the regime is more likely to change if θ is high (i.e., the status quo is weak) and/or p is high (i.e., more agents attack it). We assume that there exist $\underline{\theta}, \bar{\theta} \in \mathbf{R}$ such that $R(\underline{\theta}, 1) = R(\bar{\theta}, 0) = 0$; that is, the status quo always survives when $\theta < \underline{\theta}$ and always changes when $\theta \geq \bar{\theta}$, irrespective of the size of the aggregate attack. In contrast, if $\theta \in [\underline{\theta}, \bar{\theta}]$, the outcome depends on the size of the aggregate attack.

The payoff to an agent from attacking the regime is H if the regime changes and L otherwise, where $L < 0 < H$. Supporting the regime is a safe action with its payoff, without loss of generality, normalized to 0. The payoffs are summarized in Table 1.

We define $\gamma = -L/(H - L)$, which captures the loss of attacking the regime when the regime survives relative to the incremental benefit of a successful attack versus an

TABLE 1. Payoffs.

	Success	Failure
Attack	H	L
Support	0	0

⁴The model is based on the general regime-change model in [Vives \(2014\)](#). As explained therein, it can be reinterpreted as a model of currency crises (as in [Morris and Shin \(1998\)](#)), bank runs (as in [Rochet and Vives \(2004\)](#) or [Goldstein and Pauzner \(2005\)](#)), political revolts ([Edmond \(2013\)](#)), or sovereign debt crises ([Szkup \(2022\)](#)).

unsuccessful one. As a tie-breaking rule, we assume that an agent attacks the regime if he is indifferent between attacking and not attacking.

Agents do not observe the fundamental θ . Instead, they share a common prior that θ is uniformly distributed over $\mathbf{R}$. Each agent first observes a free private signal, $x_{i1} = \theta + \tau_1^{-1/2} \varepsilon_{i1}$, where $\varepsilon_{i1} \sim \mathcal{N}(0, 1)$ is independent and identically distributed (i.i.d.) across agents and independent of θ , and where $\tau_1 > 0$ is the precision of this signal. We can think of the first signal as capturing heterogeneous priors across agents, arising from having (i) different earlier experiences, (ii) access to different sources of information, or (iii) different interpretations of common information.

After observing the first signal, each agent i has the choice to acquire an additional signal or to make his final decision. Note that agents make their information choices after observing the initial signal, which can be interpreted as a simple form of sequential information.⁵ The additional signal costs $C > 0$ and is given by $x_{i2} = \theta + \tau^{-1/2} \varepsilon_{i2}$, where $\tau > 0$ is the signal precision, $\varepsilon_{i2} \sim \mathcal{N}(0, 1)$, i.i.d. across agents, and is independent of θ and of ε_{i1} . Note that τ is not necessarily equal to τ_1 . We also assume that C is low enough so that some agents always acquire additional information. As a tie-breaking rule, we assume that an agent makes his final decision if he is indifferent between acquiring further information and making his final decision.

The time line of the model is as follows. In the beginning, the fundamental θ is realized. Then each agent observes his initial (free) signal, after which he decides whether to acquire further information or to make his irreversible final decision. Agents who acquire additional information make their final decisions after observing their second private signal. Once all agents have made their final decisions, the regime outcome is determined and the payoffs are realized.

3. EQUILIBRIUM ANALYSIS

As is common in the literature, we focus on *threshold equilibria*, each characterized by a regime-change threshold, which we denote by θ^* , such that the regime changes if and only if $\theta \geq \theta^*$.

3.1 Individual choices

We start by analyzing the choices of agent i , who believes that the regime-change threshold is $\hat{\theta}$, for some arbitrary $\hat{\theta} \in \mathbf{R}$. Holding $\hat{\theta}$ fixed, agent i first observes his initial signal and decides whether to acquire the additional information or to make his final decision. If agent i acquires the second signal, he observes it and then makes his final decision. For each $n \in \{1, 2\}$, agent i 's posterior belief about θ after observing n signal(s) is given by $\mathcal{N}(\mu_{in}, \tau_n^{-1})$, where $\tau_n = \tau_1 + (n - 1)\tau$, $\mu_{i1} = x_{i1}$, and

$$\mu_{i2} = \frac{\tau_1 \mu_{i1} + \tau x_{i2}}{\tau_1 + \tau}.$$

⁵In Section 5, we consider an extension in which agents can sequentially acquire up to $N - 1$ signals, for arbitrary $N > 2$.

In what follows, for expositional simplicity, we drop the index i . Whenever it does not lead to confusion, we refer to μ_n as an agent's "posterior," where the subscript indicates the number of signals an agent has observed to reach such a posterior.

Optimal final decisions Fix $\hat{\theta}$ and consider an agent with posterior μ_n . The expected payoff to the agent from attacking the regime is $H \Pr(\theta \geq \hat{\theta} | \mu_n) + L \Pr(\theta < \hat{\theta} | \mu_n)$, which is increasing strictly in μ_n . Since supporting the regime yields a payoff of 0, the optimal final decision takes the form of a threshold strategy, where an agent finds it optimal to attack the regime if and only if his posterior belief is high enough.

LEMMA 1. *For each $\hat{\theta} \in \mathbf{R}$, an agent with posterior μ_n chooses to attack the regime if and only if $\mu_n \geq \mu_n^*(\hat{\theta})$, where $\mu_n^*(\hat{\theta}) = \hat{\theta} + \tau_n^{-1/2} \Phi^{-1}(\gamma)$ and Φ denotes the cumulative distribution function (CDF) of the standard normal distribution.*

The above optimal decision rule has two implications. First, when γ is low (i.e., H is high compared to L), $\mu_n^*(\hat{\theta})$ is lower than $\hat{\theta}$, as agents are willing to take more risk in order not to miss out on the high payoff associated with a successful attack. The opposite is true when γ is high. Second, the value to an agent with posterior μ_n from making the optimal final decision, $U(\mu_n)$, is given by

$$U_n(\mu_n) = \begin{cases} H \Pr(\theta \geq \hat{\theta} | \mu_n) + L \Pr(\theta < \hat{\theta} | \mu_n), & \text{if } \mu_n \geq \mu_n^* \\ 0, & \text{if } \mu_n < \mu_n^*. \end{cases} \quad (1)$$

Optimal information choice We now turn our attention to the optimal information choice of a typical agent. Let

$$B(\mu_1) \equiv \mathbf{E}[U_2(\mu_2) | \mu_1] - U_1(\mu_1) \quad (2)$$

be the benefit from acquiring the second signal to an agent with posterior μ_1 . It is straightforward to show (see Appendix A.2) that B is strictly increasing on $(-\infty, \mu_1^*(\hat{\theta}))$, is strictly decreasing on $(\mu_1^*(\hat{\theta}), \infty)$, achieves maximum at $\mu = \mu_1^*(\hat{\theta})$, and tends to 0 as $\mu_1 \rightarrow \pm\infty$. The above observations imply a simple optimal decision for information acquisition.

LEMMA 2. *For any $\hat{\theta} \in \mathbf{R}$, there exist two thresholds, $\bar{\mu}_1(\hat{\theta})$ and $\underline{\mu}_1(\hat{\theta})$, where $\bar{\mu}_1(\hat{\theta}) \geq \mu_1^*(\hat{\theta}) \geq \underline{\mu}_1(\hat{\theta})$, such that the following statements hold:*

- (i) *An agent acquires the second signal if and only if $\mu_1 \in (\underline{\mu}_1(\hat{\theta}), \bar{\mu}_1(\hat{\theta}))$.*
- (ii) *An agent chooses to attack the regime (support the regime, resp.) after observing the initial signal if and only if $\mu_1 \geq \bar{\mu}_1(\hat{\theta})$ ($\mu_1 \leq \underline{\mu}_1(\hat{\theta})$, resp.).*
- (iii) *We have $\partial \underline{\mu}_1(\hat{\theta}) / \partial \hat{\theta} = \partial \bar{\mu}_1(\hat{\theta}) / \partial \hat{\theta} = 1$.*

Lemma 2 implies that an agent makes the final decision after observing the initial signal only when one action looks sufficiently better than the other given his posterior and acquires information otherwise. It also implies that the number of signals

each agent acquires is endogenous and so is the distribution of agents' posterior beliefs. This observation differentiates our setup from those considered by the earlier literature on information acquisition (where all agents make the same information acquisition choices) and the literature on the impact of exogenous changes in information precision (which can be interpreted as providing all agents with an additional signal). The final part of Lemma 2 shows that a change in the regime-change threshold $\hat{\theta}$ simply shifts the information acquisition region without affecting its size.

3.2 Aggregate attack and equilibrium invariance

For each $\theta, \hat{\theta} \in \mathbf{R}$, define $p(\theta; \hat{\theta})$ as the aggregate attack when the fundamental is θ and the regime-change threshold is $\hat{\theta}$. Lemmas 1 and 2 imply that⁶

$$p(\theta; \hat{\theta}) = \int_{\bar{\mu}_1(\hat{\theta})}^{\infty} f(\mu_1 | \theta) d\mu_1 + \int_{\underline{\mu}_1(\hat{\theta})}^{\bar{\mu}_1(\hat{\theta})} \int_{\mu_2^*(\hat{\theta})}^{\infty} f(\mu_2 | \mu_1, \theta) f(\mu_1 | \theta) d\mu_2 d\mu_1. \quad (3)$$

An equilibrium regime-change threshold is then a solution to

$$R(\hat{\theta}, p(\hat{\theta}; \hat{\theta})) = 0, \quad (4)$$

which is commonly referred to as the *regime-change condition*.

We now state our main result, which provides a complete characterization of the unique threshold equilibrium of our game.

THEOREM 1 (Invariance). *For all $\hat{\theta} \in \mathbf{R}$, we have $p(\hat{\theta}; \hat{\theta}) = 1 - \gamma$. Therefore, the game admits a unique threshold equilibrium. The equilibrium is characterized by a regime-change threshold θ^* that uniquely satisfies $R(\theta^*, 1 - \gamma) = 0$.*

Theorem 1 establishes an invariance property for the class of regime-change global games considered in this paper, which states that the equilibrium regime-change threshold is independent of parameters governing sequential information choices such as information costs C and information precision τ . In other words, Theorem 1 indicates that information acquisition has no impact on the incidence of coordination failure when agents make their information choices after observing initial private information.

To gain intuition about the invariance result, consider, first, a simple static version of our model in which agents do not have access to the second signal. In this case, it is well known that the proportion of agents attacking the regime is equal to $1 - \gamma$. How would introducing a costly second signal affect the aggregate attack? On the one hand, the second signal leads some agents to switch from attacking to not attacking the regime. These are agents whose initial posterior beliefs belong to $[\mu_1^*, \bar{\mu}_1)$ but who receive low second signals so that their final posterior beliefs, μ_2 , lie below μ_2^* . On the other hand, agents whose initial posterior beliefs belong to $(\underline{\mu}_1, \mu_1^*)$ but who receive high second

⁶Here, $f(\mu_1 | \theta)$ is the probability density function (PDF) of μ_1 conditional on the true state taking value θ , and $f(\mu_2 | \mu_1, \theta)$ is the PDF of μ_2 conditional on the true state taking value θ and the initial posterior being μ_1 .

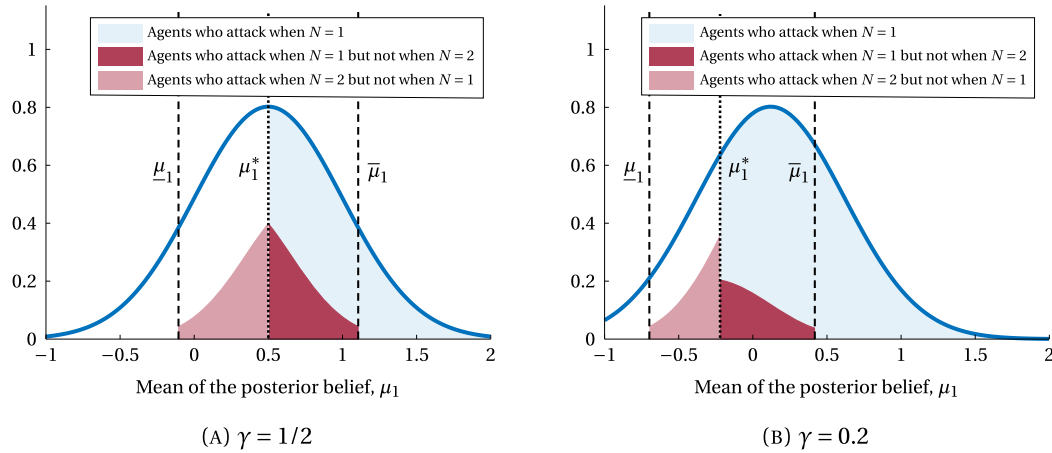


FIGURE 1. The effect of allowing agents to acquire an additional signal when agents have an improper prior. The solid blue line depicts $f(\mu_1 | \theta_1^*)$. The light blue area is equal to the proportion of agents attacking the regime in the model without additional information. The dark red area is equal to the proportion of agents who switch from attacking to not attacking when they can acquire an additional signal. Similarly, the light red area is equal to the proportion of agents who switch from not attacking to attacking when they can acquire an additional signal. Parameters: $H - L = 4$, $\tau_1 = \tau = 4$, $C = 0.025$, $R(\theta, p) = \theta + p - 1$.

signals so that $\mu_2 \geq \mu_2^*$ switch their actions in the opposite direction. The invariance result follows from the observation that, under the improper uniform prior, these two groups of agents are of equal size, so the positive and negative effects of the second signal on the size of the attack exactly cancel out.

To understand why these two groups of agents are of equal size, let N denote the total number of signals agents can observe and let θ_1^* denote the equilibrium regime change threshold when $N = 1$. Consider Figure 1, which depicts the density of posterior belief μ_1 given θ_1^* , $f(\mu_1 | \theta_1^*)$. The shaded areas below the density correspond to the proportions of agents attacking the regime in the model with a single signal (light blue area), switching from attacking to not attacking after introducing a second signal (dark red area), and vice versa (light red area), when $\gamma = 1/2$ (panel A) and $\gamma < 1/2$ (panel B).⁷ We now discuss these two cases separately.⁸

Intuition when $\gamma = 1/2$ In this case, the benefit from a successful attack is equal to the loss from an unsuccessful attack, so that $\mu_1^* = \theta_1^*$. This implies, as depicted in panel A of Figure 1, that $f(\mu_1 | \theta_1^*)$ is symmetric about μ_1^* and, thus, an agent's incentive to acquire additional information as a function of his initial posterior is also symmetric

⁷To be precise, the area below the density for $\mu_1 \geq \mu_1^*$ corresponds to the proportions of agents attacking the regime in the model with a single signal (light blue area), the area below the density between μ_1 and μ_1^* corresponds to the proportion of agents who would not attack if they could only observe a single signal but acquire an additional signal when $N = 2$, while the shaded part of this area corresponds to the proportion of agents who would switch from not attacking to attacking if they observe an additional signal (i.e., those agents whose final posterior is $\mu_2 \geq \mu_2^*$).

⁸The case $\gamma > 1/2$ is analogous to the case $\gamma < 1/2$.

(i.e., $|\underline{\mu}_1 - \mu_1^*| = |\bar{\mu}_1 - \mu_1^*|$). Therefore, equal proportions of agents with $\mu_1 < \mu_1^*$ and those with $\mu_1 > \mu_1^*$ acquire the additional signal. Since $\mu_2^* = \mu_1^*$ if $\gamma = 1/2$ (see Lemma 1), this implies that after observing the second signal, the proportion of agents who switch from attacking to supporting (dark red area) is equal to the proportion of agents who switch from supporting to attacking (light red area). Thus, when $\gamma = 1/2$, the invariance result is driven by the symmetry of $f(\mu_1|\theta_1^*)$ with respect to μ_1^* and the fact that $\mu_1^* = \mu_2^*$.

Intuition when $\gamma < 1/2$ In this case, the benefit from a successful attack exceeds the loss from an unsuccessful attack so that $\mu_1^* < \theta_1^*$. This implies that most of the distribution mass of $\mu_1|\theta_1^*$ lies to the right of μ_1^* (see panel B of Figure 1). Therefore, the majority of agents who acquire information would have attacked the regime if they had had to make their final decision after the first signal (i.e., the area below the density between $\bar{\mu}_1$ and μ_1^* is larger than the area between $\underline{\mu}_1$ and μ_1^* as depicted in panel B of Figure 1). However, agents with a posterior lower than μ_1^* who acquire the second signal are more likely to receive a signal that will shift their belief substantially upward (since $\mu_1^* < \theta_1^*$). It follows that even though there are fewer agents with a posterior lower than μ_1^* who acquire information, these agents are more likely to switch from attacking to supporting. Thus, as depicted in panel B of Figure 1, the majority of agents with an initial posterior $\mu_1 \in [\underline{\mu}_1, \mu_1^*)$ switch their final decisions if they are able to acquire the second signal. When agents have an improper prior, these two forces offset each other, implying that the overall proportion of agents attacking the regime after introducing an additional signal stays unchanged.

The role of the improper prior The invariance result depends crucially on the assumption that agents have an improper prior.⁹ To understand why, assume that agents start with a common proper prior $\theta \sim \mathcal{N}(\mu_\theta, \tau_\theta^{-1})$, where $\tau_\theta > 0$ captures the precision of the prior, and denote by $\theta_1^*(\mu_\theta)$ the equilibrium regime-change threshold as a function of μ_θ when agents observe only a single private signal. Note that the proper prior shifts the mean of the distribution of $\mu_1|\theta_1^*$ from θ_1^* to $(\tau_\theta \mu_\theta + \tau \theta_1^*)/(\tau_\theta + \tau)$. This breaks the “canceling effect” present under the improper prior: depending on the value of μ_θ , more agents switch from attacking to not attacking, or vice versa, than under an improper prior. As a result, the regime-change threshold becomes sensitive to information acquisition decisions.

Figure 2 depicts these effects in the case of low μ_θ . We see that, in this case, the presence of the proper prior shifts the mass of $\mu_1^*|\theta_1^*$ toward the interval $[\underline{\mu}_1, \mu_1^*)$, implying that the proportion of agents who switch from not attacking to attacking following the introduction of the second signal is greater than the proportion of agents who switch from attacking to not attacking. As a result, in this case, the ability to acquire an additional signal decreases the regime-change threshold, improving coordination. The opposite happens when μ_θ is high.

While the invariance result does not hold under a proper prior, it should be stressed that if the precision of the prior is low compared to the precision of private signals, the invariance result will approximately hold. In this case, changes in the cost and precision

⁹Otherwise, as discussed in Sections 5 and 6, the invariance result is remarkably robust.

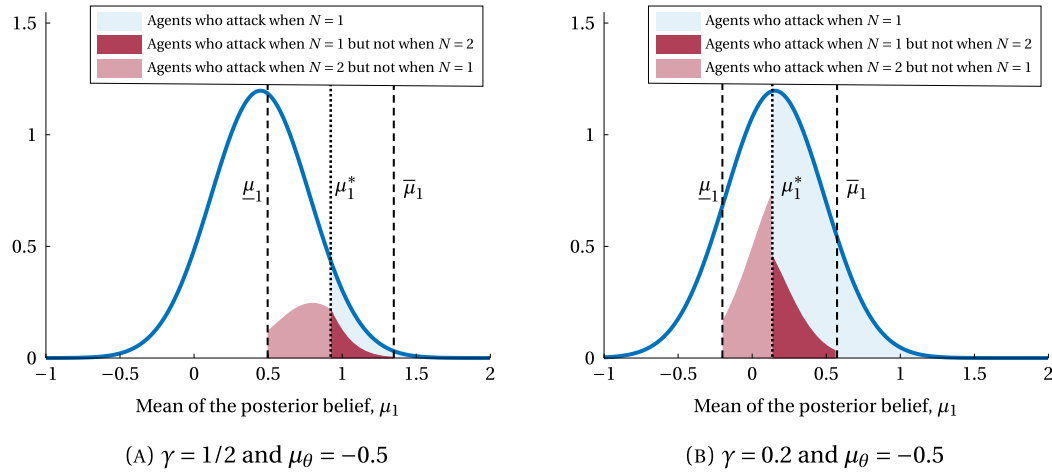


FIGURE 2. The effect of allowing agents to acquire an additional signal when agents have a proper prior. The solid line depicts $f(\mu_1|\theta_1^*)$. The light blue area is equal to the proportion of agents attacking the regime in the model without additional information (“static benchmark”). The dark red area is equal to the proportion of agents who switch from attacking to not attacking when they can acquire an additional signal. Similarly, the light red area is equal to the proportion of agents who switch from not attacking to attacking when they can acquire an additional signal. Parameters: $H - L = 4$, $\tau_1 = \tau = 4$, $\tau_\theta = 2$, $C = 0.025$, $R(\theta, p) = \theta + p - 1$.

of the information will have negligible effects. This is because, as is well established in the global games literature, a proper prior with a continuous density leads to posteriors that are close to those arising from a uniform prior when the precision of private signals is high relative to the precision of the information contained in the prior.

3.3 Discussion of the invariance result

The invariance result warrants further discussion. First, it should be clarified that the invariance result does not imply that the parameters governing information acquisition have no impact on the equilibrium. In contrast, as characterized in Proposition 1 below, changes in the cost and precision of information affect both agents’ equilibrium information acquisition strategies and the equilibrium size of aggregate attack when $\theta \neq \theta^*$. Thus, it is only the extent of coordination failure that is invariant to those parameters.

PROPOSITION 1. *The following statements are true.*

- (i) *The size of the information acquisition region measured by $\bar{\mu}_1 - \underline{\mu}_1$ is decreasing in C but is increasing in τ .*
- (ii) *A decrease in C leads $p(\theta; \theta^*)$ to decrease for all $\theta < \theta^*$ and to increase for all $\theta > \theta^*$.*

Not surprisingly, an increase in C shrinks the information acquisition region, as a higher cost deters more agents from acquiring information. Also, for every posterior

μ_1 , acquiring the second signal becomes more valuable if it is more precise (i.e., if τ is higher), and so the information acquisition region expands as τ increases. The second part of Proposition 1 follows from the observation that additional information helps agents better coordinate their actions with the regime outcome.

Second, it should be noted that the invariance result is an equilibrium result, in the sense that if agents follow non-equilibrium information acquisition strategies, then, in general, the regime-change threshold will depend on the costs and precision of signals. This point is illustrated in Section 4, where we characterize the sequential information acquisition strategy that minimizes the incidence of coordination failure.

Finally, it is worth stressing the generality of the invariance result. In particular, we derived this result in a setup that nests many of the standard global games models, including the general setup of Vives (2014). Moreover, as we discuss in Sections 5 and 6, the invariance result is remarkably robust and continues to hold in many natural extensions of our baseline model.

4. MINIMIZING COORDINATION FAILURE

In this section, we characterize the sequential information acquisition strategy that minimizes coordination failure. Our analysis is motivated by the observation that, in many cases, it is optimal from the society's point of view to decrease the incidence of coordination failure. Furthermore, identifying the most efficient way to acquire information helps us understand the inefficiency in agents' information choices. Finally, our analysis illustrates our claim that the invariance result is an equilibrium phenomenon and does not hold for arbitrary information acquisition strategies.

We consider a planner (she) whose goal is to minimize the regime-change threshold. To achieve her goal, the planner can control agents' information choices but has no control over agents' use of information.¹⁰ More formally, the planner minimizes the regime-change threshold θ^* by choosing a function $J : \mathbf{R} \rightarrow [0, 1]$, referred to as her *information policy*, which specifies the probability that she offers an agent the second private signal for each value of the initial posterior belief μ_1 . For clarity of comparison with the equilibrium, we focus on information policies that induce unique regime-change thresholds.¹¹ The policy-implied regime-change threshold is determined by agents' final decisions based on the information they observe. The planner does not take into account the cost of her information policy.¹²

¹⁰Colombo, Femminis, and Pavan (2014) consider such a planner's problem to characterize the social value of information in Gaussian-quadratic games with ex ante information acquisition.

¹¹The set of such policies includes "interval policies," where the planner provides an agent with an additional signal if and only if the agent's posterior belief belongs to an interval. An example of an interval policy is the equilibrium information acquisition policy. In the Appendix, we provide a sufficient condition that ensures that any information policy induces a threshold equilibrium (see Lemma A9).

¹²Note that the planner's problem is nontrivial since a change in her information policy leads to a change in θ^* , which then affects the optimal choice of information policy. Moreover, we do not directly restrict the type of information policy that the planner can choose. Thus, the planner faces a complex fixed-point problem.

PROPOSITION 2 (Planner's solution). *The lowest regime-change threshold the planner can achieve, denoted by $\underline{\theta}^*$, is the unique value that satisfies $R(\underline{\theta}^*, \bar{p}) = 0$, where*

$$\bar{p} = 1 - \int_{-\infty}^{\Phi^{-1}(\gamma)} \int_{-\infty}^{\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma)} \phi(z_1) \phi\left(z_2 - \sqrt{\frac{\tau_1}{\tau}} z_1\right) dz_2 dz_1$$

is the size of the aggregate attack when fundamentals take value θ^ and ϕ is the standard normal PDF. To achieve $\underline{\theta}^*$, the planner offers the second signal only to agents with posterior $\mu_1 \leq \mu_1^*(\underline{\theta}^*) \equiv \underline{\theta}^* + \tau_1^{-1/2} \Phi^{-1}(\gamma)$.*

Despite allowing for a wide array of information policies, the optimal solution is simple and intuitive: the planner offers the additional signal to an agent only if the agent would not attack the regime without the additional information. Therefore, compared with the planner's solution, agents' equilibrium information choices exhibit a *two-sided inefficiency*: too many agents with high posteriors and too few agents with low posteriors acquire information in the equilibrium.

Proposition 2 establishes that agents' information choices are inefficient. However, it is silent about the extent of this inefficiency. The next result indicates that the extent of the inefficiency is linked to the payoff parameter γ .

PROPOSITION 3. *Suppose that $R(\theta, p) = \theta + p - 1$. The difference between the equilibrium regime-change threshold, θ^* , and the lowest regime-change threshold that the planner can achieve, $\underline{\theta}^*$, is given by*

$$\theta^* - \underline{\theta}^* = \int_{\Phi^{-1}(\gamma)}^{\infty} \Phi\left(\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma) - \sqrt{\frac{\tau_1}{\tau}} z\right) \phi(z) dz.$$

Moreover, $\theta^ - \underline{\theta}^*$ is a single-peaked and symmetric function of γ , and achieves its maximum at $\gamma = 1/2$ and tends to 0 as $\gamma \rightarrow 0$ or $\gamma \rightarrow 1$.*

Proposition 3 shows that the inefficiency resulting from agents' information decisions is the smallest when γ takes extreme values. In other words, the difference $\theta^* - \underline{\theta}^*$ is small when, from the ex ante perspective, one action seems clearly preferable. Applied in the context of investment complementarities considered in Dasgupta (2007) and Szkup and Trevino (2015), this result implies that when the payoff from a successful investment is large relative to the cost of investing (or vice versa), the inefficiency resulting from information decisions is low.¹³

One may wonder if policymakers can devise subsidies to information acquisition that help to reduce coordination failure. It turns out that there is no information subsidy scheme based only on observable actions that can improve coordination. To see this point, assume that a policymaker reimburses an agent who has acquired the second signal a fraction $s \in [0, 1]$ of the information acquisition cost if the agent ends up attacking the regime or a fraction $t \in [0, 1]$ if the agent does not end up attacking. Note that

¹³In Dasgupta (2007) and Szkup and Trevino (2015), investors receive payoff $b - c$ if investment is successful (with $b > c > 0$) and $-c$ if investment is unsuccessful, with c interpreted as the cost of investment. In the context of their model, we see that $\gamma = c/b$.

the policymaker simply provides an unconditional subsidy to information acquisition if $s = t$.

PROPOSITION 4 (Subsidies). *Let $\theta^*(s, t)$ denote the equilibrium regime-change threshold with the subsidies $\{s, t\}$. For any $\{s, t\} \in [0, 1]^2$, $\theta^*(s, t) = \theta^*(0, 0)$.*

Since $\theta^*(0, 0)$ is the regime-change threshold without subsidies (as characterized in Theorem 1), the above result states that any subsidy to information acquisition that is not conditioned on agents' interim beliefs (or, equivalently, on agents' private signals) will have no impact on the extent of coordination failure.

5. BEYOND TWO SIGNALS ($N > 2$)

Above, we assumed that agents have an opportunity to acquire only one additional signal. In this section, we extend our baseline model by allowing agents to acquire up to $N - 1$ additional private signals, $N > 2$. The cost of the n th additional signal is $C_n > 0$, and the cost increases with the number of signals, that is, $C_{n+1} \geq C_n$. All additional signals have the same precision τ , and the n th signal is given by $x_{in} = \theta + \tau^{-1/2} \varepsilon_{in}$, where ε_{in} is i.i.d. across agents and independent of θ and ε_{ik} , $k \neq n$. The following theorem characterizes the threshold equilibrium in this case.

THEOREM 2 (Invariance for $N > 2$). *For all $\hat{\theta} \in \mathbf{R}$, we have $p(\hat{\theta}; \hat{\theta}) = 1 - \gamma$. Therefore, the unique threshold equilibrium is characterized by the regime-change threshold θ^* , where θ^* is the unique value satisfying $R(\theta^*, 1 - \gamma) = 0$.*

Theorem 2 shows that the invariance result extends to the case with $N > 2$ signals. As in the baseline model, this result follows from the observation that, starting with any N , introducing an additional signal leaves the volume of agents attacking the regime unchanged.

However, allowing for more signals leads to interesting new conclusions about minimizing coordination failure. As before, consider a planner whose goal is to minimize the regime-change threshold. The planner can control agents' information choices but not their use of information.¹⁴ Let $\underline{\theta}^*(N)$ denote the regime-change threshold that the planner can achieve when agents can observe up to N signals and let $\overline{p}(N)$ denote the size of the aggregate attack when fundamentals take value $\underline{\theta}^*(N)$.

PROPOSITION 5. *The following statements hold:*

- (i) *The term $\overline{p}(N)$ is strictly increasing in N and $\underline{\theta}^*(N)$ is strictly decreasing in N .*
- (ii) *We have $\lim_{N \rightarrow \infty} \overline{p}(N) = 1$ and $\lim_{N \rightarrow \infty} \underline{\theta}^*(N) = \underline{\theta}$.*

¹⁴Formally, the planner minimizes the regime-change threshold θ^* by choosing a set of functions $\mathcal{J} = \{J_n\}_{n=1}^{N-1}$, where $J_n : \mathbf{R} \rightarrow [0, 1]$ is the probability that she offers an agent the $(n + 1)$ th private signal at each possible posterior that an agent can reach after observing n signals. As before, we focus on information policies that induce unique regime-change thresholds.

Proposition 5 shows that the planner can do better with more signals, as she now can provide more signals to agents with pessimistic beliefs in order to change their minds and make them attack the regime. It also establishes that as $N \rightarrow \infty$, the proportion of agents attacking the regime converges to 1, while the regime-change threshold converges to the lower bound of the coordination region, $\underline{\theta}$. Thus, for a sufficiently large N , the planner is able to almost completely eliminate coordination failure. This is because the planner's solution guarantees that for a large enough N , at some point, almost every agent will reach a posterior higher than $\mu_n^*(\underline{\theta}^*)$, $n \in \{1, 2, \dots, N\}$, and, thus, will choose to attack the regime.¹⁵

6. FURTHER DISCUSSION AND EXTENSIONS

In this section, we briefly discuss some possible extensions of our model.

Heterogeneity

It is straightforward to show that Theorem 1 extends to the case with ex ante heterogeneous agents as in Sakovics and Steiner (2012). To see this, suppose that, as in Sakovics and Steiner (2012), there are K groups of agents ($K < \infty$) with groups indexed by k , where each group k consists of a continuum of identical agents with measure m_k , $\sum_{k=1}^K m_k = 1$. Groups may differ in terms of cost, availability, and precision of information as well as in their payoffs. Then, holding $\hat{\theta}$ fixed, we can apply the same argument as used in the proof of Theorem 1 to establish that the measure of agents in group k attacking the regime is equal to $m_k(1 - \gamma_k)$, where γ_k is the payoff parameter of group k defined analogously to γ in our benchmark model. This leads to the following result.

COROLLARY 1 (Heterogeneous agents). *For all $\hat{\theta} \in \mathbf{R}$, we have $p(\hat{\theta}; \hat{\theta}) = \sum_k m_k(1 - \gamma_k)$. Therefore, the unique threshold equilibrium in the model with heterogeneous agents is characterized by the regime-change threshold θ^* , which is the unique solution to $R(\theta^*, \sum_k m_k(1 - \gamma_k)) = 0$.*

Signals of varying precision

In the baseline model as well as in its extension to $N > 2$ signals, we have assumed for expositional ease that all additional signals have the same precision. However, by inspecting the proof of the invariance result, one can readily observe that Theorem 1 continues to hold without this assumption. Thus, our main results continue to hold when signals are of varying precision.

Learning about others' choices

In the model, we assume that agents acquire additional information directly about θ . In some situations, it might be more natural to think that agents can learn about the proportion of agents who attacked the regime in previous periods. However, as Dasgupta,

¹⁵A similar result, although in a different setting and under the equilibrium play, is established in the context of a dynamic global game model in Dasgupta, Steiner, and Stewart (2012).

Steiner, and Stewart (2012) argue, the distinction between learning about fundamentals and learning about others' final decisions in global games is, to a large extent, superficial. This is because, in a threshold equilibrium, the measure of agents attacking the regime after observing n or fewer signals is an increasing function of θ . Therefore, each signal about the fundamental can be shown to be informationally equivalent to a noisy observation of other players' past actions. Thus, our results continue to hold when agents can acquire information about the actions of others.

Multiple information sources

In the baseline model, agents' information choices are restricted to stopping or continuing information acquisition. In reality, however, people may have access to multiple information sources and, thus, can choose the source from which to acquire additional information (where information sources differ in the precision and the cost of the information they provide). It turns out that our main insights carry over to the extended version of our model that features multiple information sources. That is, the invariance result continues to hold if agents can choose the next signal to acquire from a finite set of signals, each with a different cost and precision.¹⁶

Proper Prior and Public Information

As discussed in Section 3.2, the invariance result does not hold when agents have a common proper prior. It is then natural to ask how a proper prior might influence the equilibrium behavior. To address this, we incorporate a common proper prior into the model and assume that agents initially share the belief $\theta \sim N(\mu_\theta, \tau_\theta^{-1})$ with $\tau_\theta \in (0, +\infty)$. Since the prior belief is common information, it can be interpreted as resulting from agents' observing a public signal with value μ_θ and precision τ_θ . We then ask how the mean of the prior (or, equivalently, the value of the public signal) affects agents' information choices and the equilibrium regime-change threshold.

Not surprisingly, an increase in μ_θ has a non-monotone effect on information acquisition. When μ_θ is low, few agents acquire information because most agents are confident that the regime will survive. As μ_θ increases, uncertainty about the regime outcome first intensifies because the public information becomes less indicative of the regime outcome. This increases agents' incentives for the acquisition of additional information. However, beyond a critical level, a further increase in μ_θ decreases these incentives. This is because, based on public information, most agents begin to anticipate a regime change. As a result, additional information becomes less valuable to agents since it is unlikely to change such an optimistic expectation. As such, the measure of agents who acquire information is single-peaked in μ_θ .

In contrast, an increase in μ_θ always reduces the equilibrium regime-change threshold, thereby facilitating coordination. The intuition behind this result is standard. Other factors being equal, an increase in μ_θ shifts agents' posteriors upward, which leads

¹⁶The proof of this result can be found in the working paper version (<https://drive.google.com/file/d/1IXIRbB0eQILYZchfUQxeS-4avDiVeFf3/view?usp=sharing>) of this article.

agents to assign a higher likelihood to a weak regime, encouraging each agent to attack. The nature of coordination further amplifies this effect since, now, every agent also expects that others are more likely to attack. As a consequence, an increase in μ_θ increases the size of the aggregate attack, resulting in a lower θ^* .

7. CONCLUSION

In this paper, we investigate how changes in the costs and precision of private information affect the incidence of coordination failure when agents decide sequentially how much information to acquire. To do so, we embed sequential information acquisition into a general global games model of regime change. Our main finding is the invariance result, which states that when agents have an improper uniform prior belief, the parameters governing sequential information choices have no impact on the extent of coordination failure. Our results broadly suggest that the likelihood of crises driven by coordination failure (such as bank runs, rollover crises, and sovereign debt crises) is largely unaffected by the increase in information accessibility and the decrease in its costs. While this insight is obtained in a model with an improper prior, it approximately holds when agents have a proper but diffuse prior, which, we believe, captures the heightened uncertainty typical for crisis periods well.

Our results contribute to the large theoretical literature on the effect of cheaper and more abundant information on the likelihood of financial crises. However, the results of this literature are ambiguous. For example, [Iachan and Nenov \(2015\)](#) find that whether more precise information decreases or increases the likelihood of crises depends on the sensitivity of payoffs to fundamentals. Others, such as [Szkup and Trevino \(2015\)](#) or [Vives \(2014\)](#), link the effect of more precise information to the prior belief. Finally, our paper suggests that the costs and precision of private information may leave the likelihood of coordination failure unchanged. The varying conclusions reached by these papers suggest the need for a careful empirical analysis to test the predictions of these models. This is a challenging but exciting avenue for future research.

APPENDIX A: PROOFS OF RESULTS IN SECTION 3

A.1 Proof of Lemma 1

For a fixed regime-change threshold $\hat{\theta}$, an agent with posterior μ_n chooses to attack if and only if $\Pr(\theta \geq \hat{\theta} \mid \mu_n) \geq \gamma$. Since $\theta \mid \mu_n \sim \mathcal{N}(\mu_n, \tau_n^{-1})$, this inequality can be written explicitly as $1 - \Phi(\frac{\hat{\theta} - \mu_n}{\tau_n^{-1/2}}) \geq \gamma$, from which we obtain $\mu_n \geq \hat{\theta} + \tau_n^{-1/2} \Phi^{-1}(\gamma) = \mu_n^*(\hat{\theta})$, as was to be shown.

A.2 Proof of Lemma 2

Fix an arbitrary regime-change threshold $\hat{\theta}$. It is easy to see that

$$\begin{aligned} U_n(\mu_n) &= \max\{H \Pr(\theta \geq \hat{\theta} \mid \mu_n) + L \Pr(\theta < \hat{\theta} \mid \mu_n), 0\} \\ &= \max\{\Delta \Pr(\theta \geq \hat{\theta} \mid \mu_n) + L, 0\}, \end{aligned}$$

where $\Delta \equiv H - L > 0$. An agent with posterior μ_1 should optimally acquire the second signal if and only if $B(\mu_1) > C$.¹⁷

We first argue that $B(\mu_1)$ is single-peaked and achieves its maximum at $\mu_1 = \mu_1^*(\hat{\theta})$. To this end, note that

$$B(\mu_1) = \begin{cases} \int_{\mathbf{R}} U_2(\mu_2) f(\mu_2 | \mu_1) d\mu_2, & \text{if } \mu_1 < \mu_1^*(\hat{\theta}) \\ \int_{\mathbf{R}} \max\{L, -\Delta \Pr(\theta \geq \hat{\theta} | \mu_2)\} f(\mu_2 | \mu_1) d\mu_2 - L, & \text{if } \mu_1 \geq \mu_1^*(\hat{\theta}). \end{cases}$$

Since $U_2(\mu_2)$ increases and is not constant in μ_2 while $\max\{L, -\Delta \Pr(\theta \geq \hat{\theta} | \mu_2)\}$ decreases and is not constant in μ_2 , it follows that $B(\mu_1)$ is strictly increasing on $(-\infty, \mu_1^*(\hat{\theta}))$ and strictly decreasing on $[\mu_1^*(\hat{\theta}), +\infty)$. The continuity of $B(\mu_1)$ guarantees that it achieves its maximum at $\mu_1 = \mu_1^*(\hat{\theta})$. Moreover, it is straightforward to verify that

$$\lim_{\mu_1 \rightarrow -\infty} B(\mu_1) = \lim_{\mu_1 \rightarrow +\infty} B(\mu_1) = 0.$$

According to the monotonicity and the limiting property of $B(\mu_1)$, we see that $B(\mu_1) > C$ if and only if $\mu_1 \in (\underline{\mu}_1(\hat{\theta}), \bar{\mu}_1(\hat{\theta}))$, where $B(\underline{\mu}_1(\hat{\theta})) = B(\bar{\mu}_1(\hat{\theta})) = C$, justifying part (i). Part (ii) then follows from Lemma 1 and part (i) (note that our characterization for $\underline{\mu}_1(\hat{\theta})$, $\bar{\mu}_1(\hat{\theta})$, and $\mu_1^*(\hat{\theta})$ implies that $\bar{\mu}_1(\hat{\theta}) > \mu_1^*(\hat{\theta}) > \underline{\mu}_1(\hat{\theta})$).

For part (iii), we pick an arbitrary $\varepsilon \in \mathbf{R}$, and write $U_n(\mu_n)$ and $B(\mu_1)$ explicitly as $U_n(\mu_n; \hat{\theta})$ and $B(\mu_1; \hat{\theta})$. It is easy to see that $U_n(\mu_n; \hat{\theta}) = \max\{\Delta \Pr(\theta \geq \hat{\theta} | \mu_n) + L, 0\} = \max\{\Delta \Pr(\theta \geq \hat{\theta} + \varepsilon | \mu_n + \varepsilon) + L, 0\} = U_n(\mu_n + \varepsilon; \hat{\theta} + \varepsilon)$, where $n = 1, 2$. As a result,

$$\begin{aligned} \mathbf{E}[U_2(\mu_2; \hat{\theta}) | \mu_1] &= \int_{\mu_2^*(\hat{\theta})}^{\infty} [\Delta \Pr(\theta \geq \hat{\theta} | \mu_2) + L] dF(\mu_2 | \mu_1) \\ &= \int_{\mu_2^*(\hat{\theta})}^{\infty} [\Delta \Pr(\theta \geq \hat{\theta} + \varepsilon | \mu_2 + \varepsilon) + L] dF(\mu_2 + \varepsilon | \mu_1 + \varepsilon) \\ &= \int_{\mu_2^*(\hat{\theta} + \varepsilon)}^{\infty} [\Delta \Pr(\theta \geq \hat{\theta} + \varepsilon | \mu_2) + L] dF(\mu_2 | \mu_1 + \varepsilon) \\ &= \mathbf{E}[U_2(\mu_2; \hat{\theta} + \varepsilon) | \mu_1 + \varepsilon], \end{aligned}$$

where, in the third line, we have employed the fact that $\mu_2^*(\hat{\theta} + \varepsilon) = \mu_2^*(\hat{\theta}) + \varepsilon$ (see Lemma 1). As a result, $B(\mu_1; \hat{\theta}) = \mathbf{E}[U_2(\mu_2; \hat{\theta}) | \mu_1] - U_1(\mu_1; \hat{\theta}) = \mathbf{E}[U_2(\mu_2; \hat{\theta} + \varepsilon) | \mu_1 + \varepsilon] - U_1(\mu_1 + \varepsilon; \hat{\theta} + \varepsilon) = B(\mu_1 + \varepsilon; \hat{\theta} + \varepsilon)$, which implies the claimed property of $\underline{\mu}_1(\hat{\theta})$ and $\bar{\mu}_1(\hat{\theta})$.

¹⁷Note that both $U_n(\cdot)$ and $B(\cdot)$ depend on $\hat{\theta}$.

A.3 Proof of Theorem 1

The aggregate attack in state $\hat{\theta}$ when all agents believe that $\hat{\theta}$ is the regime-change threshold is given by

$$p(\hat{\theta}; \hat{\theta}) = \int_{\mu_1^*(\hat{\theta})}^{\infty} f(\mu_1 | \hat{\theta}) d\mu_1 + \int_{\underline{\mu}_1(\hat{\theta})}^{\bar{\mu}_1(\hat{\theta})} \int_{\mu_2^*(\hat{\theta})}^{\infty} f(\mu_2 | \mu_1, \hat{\theta}) f(\mu_1 | \hat{\theta}) d\mu_2 d\mu_1.$$

By Lemma 2(iii), a change in $\hat{\theta}$ will cause each integral bound and each integrand in the expression above to shift by the same amount, leaving the value of $p(\hat{\theta}; \hat{\theta})$ unchanged.

Since $p(\hat{\theta}; \hat{\theta})$ is invariant in $\hat{\theta}$, we can treat it as a constant in $[0, 1]$ for all $\hat{\theta} \in \mathbf{R}$, and so the uniqueness of threshold equilibrium follows from the monotonicity of $R(\cdot, \cdot)$.

For the existence of a threshold equilibrium, it suffices to show that $p(\theta; \theta^*)$ is strictly increasing in θ , where $\theta^* \in \mathbf{R}$ is the unique value that satisfies $R(\theta^*, p(\theta^*, \theta^*)) = 0$. To this end, note that for any $\varepsilon \in \mathbf{R}$,

$$\begin{aligned} p(\theta + \varepsilon; \theta^*) &= \int_{\bar{\mu}_1}^{\infty} f(\mu_1 | \theta + \varepsilon) d\mu_1 + \int_{\underline{\mu}_1}^{\bar{\mu}_1} \int_{\mu_2^*}^{\infty} f(\mu_2 | \mu_1, \theta + \varepsilon) f(\mu_1 | \theta + \varepsilon) d\mu_2 d\mu_1 \\ &= \int_{\bar{\mu}_1}^{\infty} f(\mu_1 - \varepsilon | \theta) d\mu_1 + \int_{\underline{\mu}_1}^{\bar{\mu}_1} \int_{\mu_2^*}^{\infty} f(\mu_2 - \varepsilon | \mu_1 - \varepsilon, \theta) f(\mu_1 - \varepsilon | \theta) d\mu_2 d\mu_1 \\ &= \int_{\bar{\mu}_1 - \varepsilon}^{\infty} f(\mu_1 | \theta) d\mu_1 + \int_{\underline{\mu}_1 - \varepsilon}^{\bar{\mu}_1 - \varepsilon} \int_{\mu_2^* - \varepsilon}^{\infty} f(\mu_2 | \mu_1, \theta) f(\mu_1 | \theta) d\mu_2 d\mu_1 \\ &\equiv \rho(\varepsilon; \theta). \end{aligned}$$

Therefore, for each $\theta \in \mathbf{R}$, $\partial p(\hat{\theta}; \theta^*) / \partial \hat{\theta} |_{\hat{\theta}=\theta} = \partial \rho(\varepsilon; \theta) / \partial \varepsilon |_{\varepsilon=0}$. Note that

$$\left. \frac{\partial \rho(\varepsilon; \theta)}{\partial \varepsilon} \right|_{\varepsilon=0} = -\frac{\partial p(\theta; \theta^*)}{\partial \underline{\mu}_1} - \frac{\partial p(\theta; \theta^*)}{\partial \mu_2^*} - \frac{\partial p(\theta; \theta^*)}{\partial \bar{\mu}_1}.$$

It is immediate to see that the first two terms on the right-hand side of the above equality are positive. For the third term, observe that

$$-\frac{\partial p(\theta; \theta^*)}{\partial \bar{\mu}_1} = f(\bar{\mu}_1 | \theta) \left[1 - \int_{\mu_2^*}^{\infty} f(\mu_2 | \bar{\mu}_1, \theta) d\mu_2 \right] > 0,$$

which justifies the monotonicity claimed above.

We now prove the invariance result. Denote the game in our baseline model as $\mathcal{G}$ and consider an auxiliary game, $\mathcal{G}'$, identical to our baseline model except that all agents observe both private signals for free. Consider an arbitrary regime-change threshold $\hat{\theta} \in \mathbf{R}$, and denote by p and p' the aggregate attack in $\mathcal{G}$ and $\mathcal{G}'$, respectively, when $\theta = \hat{\theta}$.

Then¹⁸

$$\begin{aligned}
 p &= \int_{\bar{\mu}_1}^{\infty} f(\mu_1 | \hat{\theta}) d\mu_1 + \int_{\underline{\mu}_1}^{\bar{\mu}_1} \int_{\mu_2^*}^{\infty} f(\mu_2 | \mu_1, \hat{\theta}) f(\mu_1 | \hat{\theta}) d\mu_{2,1} \\
 &= \int_{\bar{\mu}_1}^{\infty} \int_{\mathbf{R}} f(\mu_2 | \mu_1, \hat{\theta}) f(\mu_1 | \hat{\theta}) d\mu_{2,1} + \int_{\underline{\mu}_1}^{\bar{\mu}_1} \int_{\mu_2^*}^{\infty} f(\mu_2 | \mu_1, \hat{\theta}) f(\mu_1 | \hat{\theta}) d\mu_{2,1}, \\
 p' &= \int_{\mu_2^*}^{\infty} f(\mu_2 | \hat{\theta}) d\mu_2 = \int_{\mathbf{R}} \int_{\mu_2^*}^{\infty} f(\mu_2 | \mu_1, \hat{\theta}) f(\mu_1 | \hat{\theta}) d\mu_{2,1}.
 \end{aligned}$$

It is well known that $p' = 1 - \gamma$. We aim to show $p = p'$.

To this end, we first write $p' - p$ explicitly as

$$p' - p = \int_{-\infty}^{\underline{\mu}_1} \int_{\mu_2^*}^{\infty} f(\mu_2 | \mu_1, \hat{\theta}) f(\mu_1 | \hat{\theta}) d\mu_{2,1} - \int_{\bar{\mu}_1}^{\infty} \int_{-\infty}^{\mu_2^*} f(\mu_2 | \mu_1, \hat{\theta}) f(\mu_1 | \hat{\theta}) d\mu_{2,1},$$

and so our goal is reduced to demonstrating that

$$\int_{-\infty}^{\underline{\mu}_1} \int_{\mu_2^*}^{\infty} f(\mu_2 | \mu_1, \hat{\theta}) f(\mu_1 | \hat{\theta}) d\mu_{2,1} = \int_{\bar{\mu}_1}^{\infty} \int_{-\infty}^{\mu_2^*} f(\mu_2 | \mu_1, \hat{\theta}) f(\mu_1 | \hat{\theta}) d\mu_{2,1}.$$

Note that the left-hand side of the above equality is the measure of agents who would support the regime in $\mathcal{G}$ but instead attack the regime in $\mathcal{G}'$. The right-hand side is the measure of agents who would attack the regime in $\mathcal{G}$ but instead support the regime in $\mathcal{G}'$. Their difference, thus, captures the net effect of offering everyone a free signal on the size of aggregate attack, which, as we will argue below, equals 0.

Define $S(\mu_n)$ as the expected payoff of attacking for an agent with posterior μ_n ; that is,

$$S(\mu_n) = \Delta \Pr(\theta \geq \hat{\theta} | \mu_n) + L, \quad n = 1, 2.$$

Let $T(\mu_1)$ be the value of acquiring the second signal for an agent with posterior μ_1 , so that

$$T(\mu_1) = \int_{\mu_2^*}^{\infty} [\Delta \Pr(\theta \geq \hat{\theta} | \mu_2) + L] f(\mu_2 | \mu_1) d\mu_2 - C. \quad (\text{A.1})$$

Define $Q(\mu_1) = T(\mu_1) - S(\mu_1)$.¹⁹ Employing the law of iterated expectation, we obtain

$$\begin{aligned}
 Q(\mu_1) &= \int_{\mu_2^*}^{\infty} S(\mu_2) f(\mu_2 | \mu_1) d\mu_2 - C - S(\mu_1) \\
 &= \int_{\mu_2^*}^{\infty} S(\mu_2) f(\mu_2 | \mu_1) d\mu_2 - C - \int_{\mathbf{R}} S(\mu_2) f(\mu_2 | \mu_1) d\mu_2 \\
 &= - \int_{-\infty}^{\mu_2^*} [\Delta \Pr(\theta \geq \hat{\theta} | \mu_2) + L] f(\mu_2 | \mu_1) d\mu_2 - C.
 \end{aligned} \quad (\text{A.2})$$

¹⁸For notational conciseness, we have suppressed the dependence of the thresholds on $\hat{\theta}$ and have written $d\mu_2 d\mu_1$ as $d\mu_{2,1}$.

¹⁹The dependence of these functions on $\hat{\theta}$ is suppressed for notational simplicity.

We now differentiate (A.1) with respect to μ_1 and integrate by parts to obtain

$$\begin{aligned} T'(\mu_1) &= \int_{\mu_2^*}^{\infty} S(\mu_2) \frac{\partial f(\mu_2 | \mu_1)}{\partial \mu_1} d\mu_2 \\ &= - \int_{\mu_2^*}^{\infty} S(\mu_2) \frac{\partial f(\mu_2 | \mu_1)}{\partial \mu_2} d\mu_2 \\ &= \Delta \int_{\mu_2^*}^{\infty} f(\hat{\theta} | \mu_2) f(\mu_2 | \mu_1) d\mu_2, \end{aligned} \quad (\text{A.3})$$

where the second equality is due to the fact that

$$\frac{\partial f(\mu_2 | \mu_1)}{\partial \mu_1} = - \frac{\partial f(\mu_2 | \mu_1)}{\partial \mu_2}.$$

Similarly, we can differentiate (A.2) with respect to μ_1 to obtain

$$Q'(\mu_1) = -\Delta \int_{-\infty}^{\mu_2^*} f(\hat{\theta} | \mu_2) f(\mu_2 | \mu_1) d\mu_2. \quad (\text{A.4})$$

Observe that

$$\lim_{\mu_1 \rightarrow -\infty} T(\mu_1) = \lim_{\mu_1 \rightarrow \infty} Q(\mu_1) = -C \quad (\text{A.5})$$

and that

$$T(\underline{\mu}_1) = Q(\underline{\mu}_1) = 0. \quad (\text{A.6})$$

Therefore, we arrive at

$$\begin{aligned} \frac{C}{\Delta} &= \lim_{\mu_1 \rightarrow -\infty} \frac{-T(\mu_1)}{\Delta} \\ &= \frac{\left[\int_{-\infty}^{\underline{\mu}_1} T'(\mu_1) d\mu_1 - T(\underline{\mu}_1) \right]}{\Delta} \\ &= \int_{-\infty}^{\underline{\mu}_1} \int_{\mu_2^*}^{\infty} f(\hat{\theta} | \mu_2) f(\mu_2 | \mu_1) d\mu_{2,1}, \end{aligned}$$

where the first equality is based on (A.5), the second equality employs the (extended) fundamental theorem of calculus, and the last equality uses (A.3) and (A.6). A symmetric argument with respect to $Q(\mu_1)$ gives

$$\begin{aligned} \frac{C}{\Delta} &= \lim_{\mu_1 \rightarrow \infty} \frac{-Q(\mu_1)}{\Delta} \\ &= \frac{\left[\int_{\underline{\mu}_1}^{\infty} Q'(\mu_1) d\mu_1 + Q(\underline{\mu}_1) \right]}{\Delta} = \int_{\underline{\mu}_1}^{\infty} \int_{-\infty}^{\mu_2^*} f(\hat{\theta} | \mu_2) f(\mu_2 | \mu_1) d\mu_{2,1}. \end{aligned}$$

The desired result then follows from the identity²⁰

$$f(\hat{\theta} | \mu_2) f(\mu_2 | \mu_1) = f(\mu_2 | \mu_1, \hat{\theta}) f(\mu_1 | \hat{\theta})$$

for all $(\mu_2, \mu_1, \hat{\theta}) \in \mathbf{R}^3$.

A.4 Proof of Proposition 1

Part (i) That $(\underline{\mu}_1, \bar{\mu}_1)$ shrinks as C increases follows from the observation that $B(\mu_1)$ is single-peaked (see the proof of Lemma 2). As τ increases, the second private signal becomes more (Blackwell) informative, as the information structure is Gaussian, yielding a higher value for $B(\mu_1)$ for each $\mu_1 \in \mathbf{R}$. The characterization of $\underline{\mu}_1$ and $\bar{\mu}_1$ (see the proof of Lemma 2) then implies that the size of the information acquisition region expands as τ increases.

Part (ii) Define $\bar{C} \equiv B(\mu_1^*)$. It is clear that $p(\theta; \theta^*)$ is constant in C for all $C \geq \bar{C}$, and it only remains to consider the case of $C \in [0, \bar{C})$. The equilibrium volume of agents attacking the regime in state θ , denoted by $p(\theta, C; \theta^*)$, is

$$p(\theta, C; \theta^*) = 1 - F(\bar{\mu}_1 | \theta) + \int_{\underline{\mu}_1}^{\bar{\mu}_1} [1 - F(\mu_2^* | \mu_1, \theta)] f(\mu_1 | \theta) d\mu_1.$$

Note that $\underline{\mu}_1$ and $\bar{\mu}_1$ satisfy

$$C = \int_{\mu_2^*}^{\infty} [\Delta \Pr(\theta \geq \theta^*) + L] f(\mu_2 | \underline{\mu}_1) d\mu_2 = - \int_{-\infty}^{\mu_2^*} [\Delta \Pr(\theta \geq \theta^*) + L] f(\mu_2 | \bar{\mu}_1) d\mu_2,$$

and so

$$\begin{aligned} \frac{\partial \underline{\mu}_1}{\partial C} &= \frac{1}{\Delta [1 - F(\mu_2^* | \underline{\mu}_1, \theta^*)] f(\theta^* | \underline{\mu}_1)} \\ \frac{\partial \bar{\mu}_1}{\partial C} &= - \frac{1}{\Delta F(\mu_1^* | \bar{\mu}_1, \theta^*) f(\theta^* | \bar{\mu}_1)}. \end{aligned}$$

Therefore, one deduces that

$$\begin{aligned} \frac{\Delta \partial p(\theta, C; \theta^*)}{\partial C} &= -F(\mu_2^* | \bar{\mu}_1, \theta) f(\bar{\mu}_1 | \theta) \frac{\partial \bar{\mu}_1}{\partial C} - [1 - F(\mu_2^* | \underline{\mu}_1, \theta)] f(\underline{\mu}_1 | \theta) \frac{\partial \underline{\mu}_1}{\partial C} \\ &= \frac{F(\mu_2^* | \bar{\mu}_1, \theta) f(\bar{\mu}_1 | \theta)}{F(\mu_2^* | \bar{\mu}_1, \theta^*) f(\bar{\mu}_1 | \theta^*)} - \frac{[1 - F(\mu_2^* | \underline{\mu}_1, \theta)] f(\underline{\mu}_1 | \theta)}{[1 - F(\mu_2^* | \underline{\mu}_1, \theta^*)] f(\underline{\mu}_1 | \theta^*)}. \end{aligned}$$

Therefore, $\partial p(\theta, C; \theta^*) / \partial C \geq 0$ if and only if

$$\frac{[1 - F(\mu_2^* | \underline{\mu}_1, \theta^*)] f(\underline{\mu}_1 | \theta^*)}{F(\mu_2^* | \bar{\mu}_1, \theta^*) f(\bar{\mu}_1 | \theta^*)} \geq \frac{[1 - F(\mu_2^* | \underline{\mu}_1, \theta)] f(\underline{\mu}_1 | \theta)}{F(\mu_2^* | \bar{\mu}_1, \theta) f(\bar{\mu}_1 | \theta)}$$

²⁰This identity relies on the property that $f(\hat{\theta} | \mu_1) = f(\mu_1 | \hat{\theta})$, which holds in our improper prior setup.

or, more explicitly,

$$\begin{aligned} & \frac{\Phi\left(\frac{\tau\theta^* + \tau_1\mu_1}{\sqrt{\tau}} - \frac{\tau + \tau_1}{\sqrt{\tau}}\bar{\mu}_2\right)\phi(\sqrt{\tau_1}(\mu_1 - \theta^*))}{\Phi\left(\frac{\tau + \tau_1}{\sqrt{\tau}}\bar{\mu}_2 - \frac{\tau\theta^* + \tau_1\bar{\mu}_1}{\sqrt{\tau}}\right)\phi(\sqrt{\tau_1}(\bar{\mu}_1 - \theta^*))} \\ & \geq \frac{\Phi\left(\frac{\tau\theta + \tau_1\mu_1}{\sqrt{\tau}} - \frac{\tau + \tau_1}{\sqrt{\tau}}\bar{\mu}_2\right)\phi(\sqrt{\tau_1}(\mu_1 - \theta))}{\Phi\left(\frac{\tau + \tau_1}{\sqrt{\tau}}\bar{\mu}_2 - \frac{\tau\theta + \tau_1\bar{\mu}_1}{\sqrt{\tau}}\right)\phi(\sqrt{\tau_1}(\bar{\mu}_1 - \theta))}. \end{aligned}$$

Define the right-hand side of the above inequality as function $\Gamma(\theta)$. Observe that the left-hand side of this inequality is $\Gamma(\theta^*)$. Now we show that $\Gamma(\theta)$ is strictly increasing for each $\theta \in \mathbf{R}$. To this end, note that²¹

$$\begin{aligned} \Gamma'(\theta) & \stackrel{\text{sgn}}{=} \sqrt{\tau} \left[\Phi\left(\frac{\tau + \tau_1}{\sqrt{\tau}}\mu_2^* - \frac{\tau\theta + \tau_1\bar{\mu}_1}{\sqrt{\tau}}\right) \phi\left(\frac{\tau\theta + \tau_1\mu_1}{\sqrt{\tau}} - \frac{\tau + \tau_1}{\sqrt{\tau}}\mu_2^*\right) \right. \\ & \quad \left. + \Phi\left(\frac{\tau\theta + \tau_1\mu_1}{\sqrt{\tau}} - \frac{\tau + \tau_1}{\sqrt{\tau}}\mu_2^*\right) \phi\left(\frac{\tau + \tau_1}{\sqrt{\tau}}\mu_2^* - \frac{\tau\theta + \tau_1\bar{\mu}_1}{\sqrt{\tau}}\right) \right] \\ & \quad - \Phi\left(\frac{\tau + \tau_1}{\sqrt{\tau}}\mu_2^* - \frac{\tau\theta + \tau_1\bar{\mu}_1}{\sqrt{\tau}}\right) \Phi\left(\frac{\tau\theta + \tau_1\mu_1}{\sqrt{\tau}} - \frac{\tau + \tau_1}{\sqrt{\tau}}\mu_2^*\right) \tau_1(\bar{\mu}_1 - \mu_1) \\ & \stackrel{\text{sgn}}{=} \frac{\phi\left(\frac{\tau\theta + \tau_1\mu_1}{\sqrt{\tau}} - \frac{\tau + \tau_1}{\sqrt{\tau}}\mu_2^*\right)}{\Phi\left(\frac{\tau\theta + \tau_1\mu_1}{\sqrt{\tau}} - \frac{\tau + \tau_1}{\sqrt{\tau}}\mu_2^*\right)} + \frac{\phi\left(\frac{\tau + \tau_1}{\sqrt{\tau}}\mu_2^* - \frac{\tau\theta + \tau_1\bar{\mu}_1}{\sqrt{\tau}}\right)}{\Phi\left(\frac{\tau + \tau_1}{\sqrt{\tau}}\mu_2^* - \frac{\tau\theta + \tau_1\bar{\mu}_1}{\sqrt{\tau}}\right)} \\ & \quad - \frac{\tau_1}{\sqrt{\tau}}(\bar{\mu}_1 - \mu_1). \end{aligned} \tag{A.7}$$

Define $A = \frac{\tau_1}{\sqrt{\tau}}(\mu_1 - \bar{\mu}_1)$ and $g(t) = \phi(t)/\Phi(t)$, and let $y(\theta) = \frac{\tau\theta + \tau_1\mu_1}{\sqrt{\tau}} - \frac{\tau + \tau_1}{\sqrt{\tau}}\mu_2^*$. Then we can rewrite (A.7) as

$$\Gamma'(\theta) \stackrel{\text{sgn}}{=} g(y(\theta)) + g(A - y(\theta)) + A.$$

To proceed, we need the following lemma.

LEMMA A1. *We have $g(t) + t > 0$ for all $t \in \mathbf{R}$.*

PROOF OF LEMMA A1. This is equivalent to $\phi(t) + t\Phi(t) > 0$ for all $t \in \mathbf{R}$. Let $h(t) = \phi(t) + t\Phi(t)$. It is straightforward to show that $h'(t) = \Phi(t) > 0$ and $\lim_{t \rightarrow -\infty} h(t) = 0$, justifying the claimed property. $\square$

Note that $g(-t)$ is the *hazard rate* of the standard normal distribution at t , which is well known to be (strictly) convex, and so $g(t)$ is (strictly) convex as well. Thus, by the

²¹We define $\stackrel{\text{sgn}}{=}$ as the binary relation on $\mathbf{R}$ such that $a \stackrel{\text{sgn}}{=} b$ if and only if a and b have the same sign.

convexity of $g(t)$ and Lemma A1, we have

$$g(y(\theta)) + g(A - y(\theta)) + A \geq 2g\left(\frac{A}{2}\right) + A \stackrel{\text{sgn}}{=} g\left(\frac{A}{2}\right) + \frac{A}{2} > 0,$$

which implies that $\Gamma'(\theta) > 0$, and so $\partial p(\theta, C; \theta^*) / \partial C \stackrel{\text{sgn}}{=} \theta^* - \theta$, as was to be shown.

APPENDIX B: PROOFS OF RESULTS IN SECTION 4

B.1 Proof of Proposition 2

We start with a formulation of the planner's problem. For an arbitrarily given $\hat{\theta} \in \mathbf{R}$, denote by $p(\theta; \hat{\theta}, J)$ the aggregate attack when (i) the true state is θ , (ii) everyone takes $\hat{\theta}$ as the regime-change threshold, and (iii) the planner's policy is J . Explicitly, we have

$$\begin{aligned} p(\theta; \hat{\theta}, J) &= \int_{\mu_1^*(\hat{\theta})}^{\infty} f(\mu_1 | \theta) [1 - J(\mu_1)] d\mu_1 \\ &\quad + \int_{\mathbf{R}} \int_{\mu_2^*(\hat{\theta})}^{\infty} f(\mu_2 | \mu_1, \theta) f(\mu_1 | \theta) J(\mu_1) d\mu_{2,1}, \end{aligned} \quad (\text{A.8})$$

where $\mu_n^*(\hat{\theta}) = \hat{\theta} + \tau_n^{-1/2} \Phi^{-1}(\gamma)$, $n = 1, 2$.

An information policy J is said to induce a threshold equilibrium with regime-change threshold $\hat{\theta}$ if $\text{sgn}[R(\theta, p(\theta; \hat{\theta}, J))] = \text{sgn}(\theta - \hat{\theta})$ for all $\theta \in \mathbf{R}$. The planner's problem can, thus, be formally stated as

$$\min_J \hat{\theta}, \quad \text{subject to } \text{sgn}[R(\theta; p(\theta; \hat{\theta}, J))] = \text{sgn}(\theta - \hat{\theta}) \text{ for all } \theta \in \mathbf{R}.$$

The relaxed problem of the planner inherits the objective function but has a looser constraint:

$$\min_J \hat{\theta}, \quad \text{subject to } R(\hat{\theta}, p(\hat{\theta}; \hat{\theta}, J)) = 0.$$

An information policy J is said to admit $\hat{\theta}$ as a quasi-regime-change threshold if $R(\hat{\theta}, p(\hat{\theta}; \hat{\theta}, J)) = 0$.

LEMMA A2. *For every information policy J , there exists some $\hat{\theta} \in \mathbf{R}$ such that J admits $\hat{\theta}$ as a quasi-regime-change threshold.*

PROOF OF LEMMA A2. Fix an arbitrary information policy J , and for each $\tilde{\theta} \in [\underline{\theta}, \bar{\theta}]$, define $\psi_J(\tilde{\theta})$ as the unique solution to the equation about θ : $R(\theta, p(\tilde{\theta}; \tilde{\theta}, J)) = 0$ (i.e., $R(\psi_J(\tilde{\theta}), p(\tilde{\theta}; \tilde{\theta}, J)) = 0$). Thus, $\psi_J(\tilde{\theta})$ is the fundamental level required to trigger a regime change when the aggregate attack equals $p(\tilde{\theta}; \tilde{\theta}, J)$. It is clear that $\psi_J(\tilde{\theta})$ is continuous, and so by Brouwer's fixed point theorem, there is $\hat{\theta} \in [\underline{\theta}, \bar{\theta}]$ such that $\psi_J(\hat{\theta}) = \hat{\theta}$, as was to be shown. $\square$

For any information policy J and any $\varepsilon \in \mathbf{R}$, the ε -shift of J , denoted by J^ε , is the information policy such that $J^\varepsilon(\mu_1) = J(\mu_1 - \varepsilon)$ for all $\mu_1 \in \mathbf{R}$.

LEMMA A3. For each information policy J , each $\hat{\theta} \in \mathbf{R}$, and $\varepsilon \in \mathbf{R}$, we have

$$p(\hat{\theta}; \hat{\theta}, J) = p(\hat{\theta} + \varepsilon; \hat{\theta} + \varepsilon, J^\varepsilon).$$

PROOF OF LEMMA A3. By definition,

$$\begin{aligned} p(\hat{\theta} + \varepsilon; \hat{\theta} + \varepsilon, J^\varepsilon) &= \int_{\mu_1^*(\hat{\theta} + \varepsilon)}^{\infty} f(\mu_1 | \theta + \varepsilon) [1 - J^\varepsilon(\mu_1)] d\mu_1 \\ &\quad + \int_{\mathbf{R}} \int_{\mu_2^*(\hat{\theta} + \varepsilon)}^{\infty} f(\mu_2 | \mu_1, \theta + \varepsilon) f(\mu_1 | \theta + \varepsilon) J^\varepsilon(\mu_1) d\mu_{2,1}. \end{aligned}$$

Observe that

$$\begin{aligned} \int_{\mu_1^*(\hat{\theta} + \varepsilon)}^{\infty} f(\mu_1 | \hat{\theta} + \varepsilon) [1 - J^\varepsilon(\mu_1)] d\mu_1 &= \int_{\mu_1^*(\hat{\theta}) + \varepsilon}^{\infty} f(\mu_1 | \hat{\theta} + \varepsilon) [1 - J(\mu_1 - \varepsilon)] d\mu_1 \\ &= \int_{\mu_1^*(\hat{\theta})}^{\infty} f(\mu_1 + \varepsilon | \hat{\theta} + \varepsilon) [1 - J(\mu_1)] d\mu_1 \\ &= \int_{\mu_1^*(\hat{\theta})}^{\infty} f(\mu_1 | \hat{\theta}) [1 - J(\mu_1)] d\mu_1 \end{aligned}$$

and that

$$\begin{aligned} &\int_{\mathbf{R}} \int_{\mu_2^*(\hat{\theta} + \varepsilon)}^{\infty} f(\mu_2 | \mu_1, \hat{\theta} + \varepsilon) f(\mu_1 | \hat{\theta} + \varepsilon) J^\varepsilon(\mu_1) d\mu_{2,1} \\ &= \int_{\mathbf{R}} \int_{\mu_2^*(\hat{\theta}) + \varepsilon}^{\infty} f(\mu_2 | \mu_1, \hat{\theta} + \varepsilon) f(\mu_1 | \hat{\theta} + \varepsilon) J(\mu_1 - \varepsilon) d\mu_{2,1} \\ &= \int_{\mathbf{R}} \int_{\mu_2^*(\hat{\theta})}^{\infty} f(\mu_2 + \varepsilon | \mu_1 + \varepsilon, \hat{\theta} + \varepsilon) f(\mu_1 + \varepsilon | \hat{\theta} + \varepsilon) J(\mu_1) d\mu_{2,1} \\ &= \int_{\mathbf{R}} \int_{\mu_2^*(\hat{\theta})}^{\infty} f(\mu_2 | \mu_1, \hat{\theta}) f(\mu_1 | \hat{\theta}) J(\mu_1) d\mu_{2,1}. \end{aligned}$$

Comparing the sum of the last line of each of the two arrays of equalities above with $p(\hat{\theta}; \hat{\theta}, J)$ yields the desired result. $\square$

LEMMA A4 (Optimality criterion). An information policy J admitting $\hat{\theta}$ as a quasi-regime-change threshold solves the relaxed problem if there does not exist any information policy J' such that $p(\hat{\theta}; \hat{\theta}, J') > p(\hat{\theta}; \hat{\theta}, J)$.

PROOF OF LEMMA A4. Consider the contrapositive of the statement. If there exists some J' such that $p(\hat{\theta}; \hat{\theta}, J') > p(\hat{\theta}; \hat{\theta}, J)$, then by the monotonicity of $R(\cdot, \cdot)$, we must have $R(\hat{\theta}, p(\hat{\theta}; \hat{\theta}, J')) > R(\hat{\theta}, p(\hat{\theta}; \hat{\theta}, J)) = 0$, and so there must exist some $\varepsilon > 0$ such that

$$R(\hat{\theta} - \varepsilon, p(\hat{\theta}; \hat{\theta}, J')) = 0.$$

Note that by Lemma A3, $p(\hat{\theta}; \hat{\theta}, J') = p(\hat{\theta} - \varepsilon; \hat{\theta} - \varepsilon, (J')^{-\varepsilon})$, and so $R(\hat{\theta} - \varepsilon, p(\hat{\theta} - \varepsilon; \hat{\theta} - \varepsilon, (J')^{-\varepsilon})) = 0$; that is, the information policy $(J')^{-\varepsilon}$ (i.e., the information policy obtained by shifting J' by $-\varepsilon$) admits a lower quasi-regime-change threshold $\hat{\theta} - \varepsilon$ and, hence, J does not solve the relaxed problem. $\square$

LEMMA A5. *There is a unique (up to a Lebesgue null set in $\mathbf{R}$) information policy $J_{\hat{\theta}}$ such that $J_{\hat{\theta}} = \mathbf{1}_{(-\infty, \mu_1^*(\hat{\theta}))}$, J admits $\hat{\theta}$ as a quasi-regime-change threshold, and $\mu_1^*(\hat{\theta}) = \hat{\theta} + \tau_1^{-1/2} \Phi^{-1}(\gamma)$.*

PROOF OF LEMMA A5. Because for any real number ε , $J_{\hat{\theta}+\varepsilon} = J_{\hat{\theta}}^\varepsilon$, by Lemma A3, we see that $p(\hat{\theta}; \hat{\theta}, J_{\hat{\theta}})$ is constant in $\hat{\theta}$. Denote this constant by $\bar{p}$. Thus, the equation about θ' , $R(\theta', \bar{p}) = 0$, has a unique solution $\underline{\theta}^* \in [\underline{\theta}, \bar{\theta}]$. Therefore, $J_{\underline{\theta}^*}$ is the unique (up to a Lebesgue null set in $\mathbf{R}$) information policy that satisfies the conditions stated in Lemma A5. $\square$

LEMMA A6. *The information policy $J_{\underline{\theta}^*} = \mathbf{1}_{(-\infty, \mu_1^*(\underline{\theta}^*))}$, where $\underline{\theta}^*$ is defined in the proof of Lemma A5, solves the relaxed problem.*

PROOF OF LEMMA A6. It suffices to show that $J_{\underline{\theta}^*}$ satisfies the optimality criterion established in Lemma A4. This is obvious, as any deviation from $J_{\underline{\theta}^*}$ involves either offering an extra signal to agents whose posteriors will induce an attack without further information or withholding an extra signal to agents who will not attack without further information. Thus, for any information policy J' , we must have $p(\underline{\theta}^*; \underline{\theta}^*, J') \leq p(\underline{\theta}^*; \underline{\theta}^*, J_{\underline{\theta}^*})$, as was to be shown. $\square$

Notice that the information policy $J_{\underline{\theta}^*}$ is an interval policy. An interval policy solves the relaxed problem if and only if it solves the planner's problem, as is formally stated and proven below.

LEMMA A7. *An interval policy induces a threshold equilibrium with regime-change threshold $\hat{\theta}$ if and only if it admits $\hat{\theta}$ as a quasi-regime-change threshold.*

PROOF OF LEMMA A7. The necessity is obvious. For sufficiency, let $J = \mathbf{1}_{(\underline{\nu}_1, \bar{\nu}_1)}$ be an arbitrary interval policy and let $\hat{\theta}$ be the quasi-regime-change threshold it admits. By (A.8), we have

$$p(\theta; \hat{\theta}, J) = \int_{\max\{\bar{\nu}_1, \mu_1^*(\hat{\theta})\}}^{\infty} f(\mu_1 | \theta) d\mu_1 + \int_{\mathbf{R}} \int_{\mu_2^*(\hat{\theta})}^{\infty} f(\mu_2 | \mu_1, \theta) f(\mu_1 | \theta) J(\mu_1) d\mu_2 d\mu_1.$$

Employing an argument similar to the proof of Theorem 1, we can show that $p(\theta; \hat{\theta}, J)$ is strictly increasing in θ , which yields the desired result. $\square$

We have shown that $J_{\underline{\theta}^*}$ solves the planner's problem. To complete the proof, it only remains to show that $\bar{p}$ has the form given in Proposition 2. To this end, observe that

$$\bar{p} = 1 - \int_{-\infty}^{\mu_1^*(\underline{\theta}^*)} \int_{-\infty}^{\mu_2^*(\underline{\theta}^*)} f(\mu_2 | \mu_1, \underline{\theta}^*) f(\mu_1 | \underline{\theta}^*) d\mu_2 d\mu_1$$

$$= 1 - \int_{-\infty}^{\Phi^{-1}(\gamma)} \int_{-\infty}^{\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma)} \phi(\nu_1) \phi\left(\nu_2 - \sqrt{\frac{\tau_1}{\tau}} \nu_1\right) d\nu_{2,1},$$

where, from the first line to the second, we have employed the definition of $\mu_n^*(\underline{\theta}^*)$, the fact that

$$\mu_2 \mid (\mu_1, \theta) \sim \mathcal{N}\left(\frac{\tau\theta + \tau_1\mu_1}{\tau + \tau_1}, \frac{\tau}{(\tau + \tau_1)^2}\right), \quad \mu_1 \mid \theta \sim \mathcal{N}(\theta, \tau_1^{-1}),$$

and have conducted the change-of-variables

$$\nu_1 = \sqrt{\tau_1}(\mu_1 - \underline{\theta}^*) \quad \text{and} \quad \nu_2 = \sqrt{\frac{\tau_2}{\tau}}(\mu_2 - \underline{\theta}^*).$$

Our proof is now complete.

B.2 Proof of Proposition 3

According to Theorem 1 and Proposition 2, we have $\theta^* = \gamma$ and

$$\begin{aligned} \underline{\theta}^* &= 1 - \bar{p} \\ &= \int_{-\infty}^{\Phi^{-1}(\gamma)} \int_{-\infty}^{\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma)} \phi(\nu_1) \phi\left(\nu_2 - \sqrt{\frac{\tau_1}{\tau}} \nu_1\right) d\nu_{2,1} \\ &= \int_{-\infty}^{\Phi^{-1}(\gamma)} \Phi\left(\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma) - \sqrt{\frac{\tau_1}{\tau}} \nu_1\right) \phi(\nu_1) d\nu_1. \end{aligned}$$

Note that

$$\int_{\mathbf{R}} \Phi\left(\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma) - \sqrt{\frac{\tau_1}{\tau}} \nu_1\right) \phi(\nu_1) d\nu_1 = \Phi\left(\frac{\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma)}{\sqrt{1 + \frac{\tau_1}{\tau}}}\right) = \gamma,$$

from which we obtain

$$\begin{aligned} \theta^* - \underline{\theta}^* &= \gamma - \int_{-\infty}^{\Phi^{-1}(\gamma)} \Phi\left(\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma) - \sqrt{\frac{\tau_1}{\tau}} \nu_1\right) \phi(\nu_1) d\nu_1 \\ &= \int_{\Phi^{-1}(\gamma)}^{\infty} \Phi\left(\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma) - \sqrt{\frac{\tau_1}{\tau}} \nu_1\right) \phi(\nu_1) d\nu_1. \end{aligned}$$

Now we write θ^* and $\underline{\theta}^*$ as θ_γ^* and $\underline{\theta}_\gamma^*$, respectively, to make explicit their dependence on γ . Note that

$$\begin{aligned} \theta_{1-\gamma}^* - \underline{\theta}_{1-\gamma}^* &= \int_{\Phi^{-1}(1-\gamma)}^{\infty} \Phi\left(\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(1-\gamma) - \sqrt{\frac{\tau_1}{\tau}} \nu_1\right) \phi(\nu_1) d\nu_1 \\ &= \int_{-\Phi^{-1}(\gamma)}^{\infty} \Phi\left(-\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma) - \sqrt{\frac{\tau_1}{\tau}} \nu_1\right) \phi(\nu_1) d\nu_1 \end{aligned}$$

$$\begin{aligned}
&= \int_{-\infty}^{\Phi^{-1}(\gamma)} \left[1 - \Phi \left(\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma) - \sqrt{\frac{\tau_1}{\tau}} \nu_1 \right) \right] \phi(\nu_1) d\nu_1 \\
&= \gamma - \int_{-\infty}^{\Phi^{-1}(\gamma)} \Phi \left(\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma) - \sqrt{\frac{\tau_1}{\tau}} \nu_1 \right) \phi(\nu_1) d\nu_1 \\
&= \theta_\gamma^* - \underline{\theta}_\gamma^*,
\end{aligned}$$

establishing the claimed symmetry. Finally, observe that

$$\frac{\partial(\theta_\gamma^* - \underline{\theta}_\gamma^*)}{\partial \gamma} = 1 - 2\Phi \left(\frac{\sqrt{\tau_2} - \sqrt{\tau_1}}{\sqrt{\tau}} \Phi^{-1}(\gamma) \right),$$

which implies that $\theta_\gamma^* - \underline{\theta}_\gamma^*$ is strictly increasing when $\gamma \in (0, 1/2)$ and is strictly decreasing when $\gamma \in (1/2, 1)$.

B.3 Proof of Proposition 4

We state a preliminary result that we will rely on below.

LEMMA A8. *Let (s, t) be the conditional subsidies the policymaker offers. Then agents attack the regime following acquisition of the second signal if and only if their posterior after observing the second signal is larger than $\mu_2^*(\hat{\gamma})$, where*

$$\mu_2^*(\hat{\gamma}) = \begin{cases} -\infty & \text{if } \hat{\gamma} \leq 0 \\ \theta^* + \tau_2^{-1/2} \Phi^{-1}(\hat{\gamma}) & \text{if } \hat{\gamma} \in (0, 1) \\ +\infty & \text{if } \hat{\gamma} \geq 1 \end{cases} \quad (\text{A.9})$$

and $\hat{\gamma} = \gamma + (t - s) \frac{C}{H - L}$.

PROOF OF LEMMA A8. This result follows immediately from solving the indifference condition when agents receive conditional subsidies (s, t) after acquiring the second signal. $\square$

PROOF OF PROPOSITION 4. When $\hat{\gamma} \leq 0$ or $\hat{\gamma} \geq 1$, the result follows from the observation that no agent will acquire information. Thus, in this case, the model becomes a static one with an improper prior, in which the aggregate attack is well known to be equal to $1 - \gamma$ in equilibrium (see, for example, [Morris and Shin \(2004\)](#)).

Thus, we focus on the case of $\hat{\gamma} \in (0, 1)$. Consider, first, the indifference condition that determines $\underline{\mu}_1$, which (in the presence of conditional subsidies) is given by²²

$$\int_{\mu_2^*}^{\infty} \left\{ H \left[1 - \Phi \left(\frac{\theta^* - \mu_2}{\tau_2^{-1/2}} \right) \right] + L \Phi \left(\frac{\theta^* - \mu_2}{\tau_2^{-1/2}} \right) \right\} \sqrt{\frac{\tau_1 \tau_2}{\tau}} \phi \left(\frac{\mu_2 - \underline{\mu}_1}{\sqrt{\frac{\tau}{\tau_1 \tau_2}}} \right) d\mu_2$$

²²For notational simplicity, we have suppressed the dependence of $\underline{\mu}_1$, $\bar{\mu}_1$, and μ_2^* on $\hat{\gamma}$.

$$+ sC \int_{\mu_2^*}^{\infty} \sqrt{\frac{\tau_1 \tau_2}{\tau}} \phi \left(\frac{\mu_2 - \underline{\mu}_1}{\sqrt{\frac{\tau}{\tau_1 \tau_2}}} \right) d\mu_2 + tC \int_{-\infty}^{\mu_2^*} \sqrt{\frac{\tau_1 \tau_2}{\tau}} \phi \left(\frac{\mu_2 - \underline{\mu}_1}{\sqrt{\frac{\tau}{\tau_1 \tau_2}}} \right) d\mu_2 = C.$$

Dividing by $H - L$, performing a change-of-variables $z = \sqrt{\tau_2}(\mu_2 - \theta^*)$, integrating by parts, and simplifying, we obtain

$$(1 - \gamma) - \int_{\Phi^{-1}(\hat{\gamma})}^{\infty} \phi(z) \Phi \left(\sqrt{\frac{\tau_1}{\tau}} (z - \underline{a}) \right) dz = \frac{C(1 - s)}{H - L}, \quad (\text{A.10})$$

where $\underline{a} = \sqrt{\tau_2}(\underline{\mu}_1 - \theta^*)$. Next, consider the indifference condition that determines $\bar{\mu}_1$, which (in the presence of conditional subsidies) is given by

$$\begin{aligned} & \int_{-\infty}^{\mu_2^*} \left\{ H \left[1 - \Phi \left(\frac{\theta^* - \mu_2}{\tau_2^{-1/2}} \right) \right] + L \Phi \left(\frac{\theta^* - \mu_2}{\tau_2^{-1/2}} \right) \right\} \sqrt{\frac{\tau_1 \tau_2}{\tau}} \phi \left(\frac{\mu_2 - \bar{\mu}_1}{\sqrt{\frac{\tau}{\tau_1 \tau_2}}} \right) d\mu_2 \\ & + sC \int_{\mu_2^*}^{\infty} \sqrt{\frac{\tau_1 \tau_2}{\tau}} \phi \left(\frac{\mu_2 - \bar{\mu}_1}{\sqrt{\frac{\tau}{\tau_1 \tau_2}}} \right) d\mu_2 + tC \int_{-\infty}^{\mu_2^*} \sqrt{\frac{\tau_1 \tau_2}{\tau}} \phi \left(\frac{\mu_2 - \bar{\mu}_1}{\sqrt{\frac{\tau}{\tau_1 \tau_2}}} \right) d\mu_2 = C. \end{aligned}$$

Following similar steps as in the case of the equation determining $\underline{\mu}_1$, we arrive at

$$\int_{-\infty}^{\Phi^{-1}(\hat{\gamma})} \phi(z) \Phi \left(\sqrt{\frac{\tau_1}{\tau}} (z - \bar{a}) \right) dz = \frac{C(1 - s)}{H - L}, \quad (\text{A.11})$$

where $\bar{a} \equiv \sqrt{\tau_2}(\bar{\mu}_1 - \theta^*)$. Equations (A.10) and (A.11) imply that

$$\int_{-\infty}^{\Phi^{-1}(\hat{\gamma})} \phi(z) \Phi \left(\sqrt{\frac{\tau_1}{\tau}} (z - \bar{a}) \right) dz + \int_{\Phi^{-1}(\hat{\gamma})}^{\infty} \phi(z) \Phi \left(\sqrt{\frac{\tau_1}{\tau}} (z - \underline{a}) \right) dz = 1 - \gamma. \quad (\text{A.12})$$

Finally, note that the proportion of agents attacking the regime is given by

$$p(\theta^*; \theta^*) = \int_{\bar{\mu}_1}^{\infty} f(\mu_1 | \theta^*) d\mu_1 + \int_{\underline{\mu}_1}^{\bar{\mu}_1} \int_{\mu_2^*}^{\infty} f(\mu_2 | \mu_1, \theta^*) d\mu_2 d\mu_1.$$

Using the functional forms of all conditional densities, observing that $f(\mu_1 | \theta) = f(\theta | \mu_1)$ and $f(\mu_2 | \mu_1, \theta) f(\mu_1 | \theta) = f(\theta | \mu_2) f(\mu_2 | \mu_1)$, and performing change-of-variables $z = \sqrt{\tau_2}(\mu_2 - \theta^*)$, we obtain

$$p(\theta^*; \theta^*) = \int_{-\infty}^{\Phi^{-1}(\hat{\gamma})} \phi(z) \Phi \left(\sqrt{\frac{\tau_1}{\tau}} (z - \bar{a}) \right) dz + \int_{\Phi^{-1}(\hat{\gamma})}^{\infty} \phi(z) \Phi \left(\sqrt{\frac{\tau_1}{\tau}} (z - \underline{a}) \right) dz, \quad (\text{A.13})$$

where $\bar{a}$ and $\underline{a}$ are defined as above. Therefore, (A.12) and (A.13) imply that $p(\theta^*; \theta^*) = 1 - \gamma$, which is independent of the subsidies. This completes the proof. $\square$

B.4 A sufficient condition that guarantees the existence of threshold equilibria

LEMMA A9. *Every information policy induces a threshold equilibrium if*

$$\frac{\underline{R}_\theta}{\underline{R}_p} > \sqrt{\frac{6 + 4\sqrt{2}}{\pi}} \max\{\sqrt{\tau_1}, \sqrt{\tau}\},$$

where $\underline{R}_\theta \equiv \inf_{(\theta, p) \in [\underline{\theta}, \bar{\theta}] \times [0, 1]} R_\theta(\theta, p)$ and $\bar{R}_p \equiv \sup_{(\theta, p) \in [\underline{\theta}, \bar{\theta}] \times [0, 1]} R_p(\theta, p)$.

PROOF OF LEMMA A9. Fix an arbitrary information policy J and let $\tilde{\theta}$ satisfy $R(\tilde{\theta}, p(\tilde{\theta}; \tilde{\theta}, J)) = 0$ (the existence of such $\tilde{\theta}$ has been established in Lemma A2). For each $\theta \in \mathbf{R}$, employing the identity that for all $(\mu_2, \mu_1, \theta) \in \mathbf{R}^3$,

$$f(\mu_2 | \mu_1, \theta) f(\mu_1 | \theta) = f(\theta | \mu_2) f(\mu_2 | \mu_1),$$

we have²³

$$\begin{aligned} p(\theta; \tilde{\theta}, J) &= \int_{\mu_1^*}^{\infty} f(\mu_1 | \theta) [1 - J(\mu_1)] d\mu_1 + \int_{\mathbf{R}} \int_{\mu_2^*}^{\infty} f(\mu_2 | \mu_1, \theta) f(\mu_1 | \theta) J(\mu_1) d\mu_{2,1} \\ &= \underbrace{\int_{\mu_1^*}^{\infty} f(\mu_1 | \theta) [1 - J(\mu_1)] d\mu_1}_{\equiv p_1(\theta; \tilde{\theta}, J)} + \underbrace{\int_{\mathbf{R}} \int_{\mu_2^*}^{\infty} f(\theta | \mu_2) f(\mu_2 | \mu_1) J(\mu_1) d\mu_{2,1}}_{\equiv p_2(\theta; \tilde{\theta}, J)}. \end{aligned}$$

Note that

$$\begin{aligned} \left| \frac{\partial p_1(\theta; \tilde{\theta}, J)}{\partial \theta} \right| &= \left| \int_{\mu_1^*}^{\infty} \frac{\partial f(\theta | \mu_1)}{\partial \theta} [1 - J(\mu_1)] d\mu_1 \right| \\ &\leq \int_{\mu_1^*}^{\infty} \left| \frac{\partial f(\theta | \mu_1)}{\partial \theta} \right| d\mu_1 \\ &= \int_{\mu_1^*}^{\infty} \left| -\frac{\partial f(\theta | \mu_1)}{\partial \mu_1} \right| d\mu_1 \leq 2 \max_{\mu_1 \in \mathbf{R}} f(\mu_1 | \theta) \leq \sqrt{\frac{2}{\pi}} \max\{\sqrt{\tau_1}, \sqrt{\tau}\} \end{aligned}$$

and that

$$\begin{aligned} \left| \frac{\partial p_2(\theta; \tilde{\theta}, J)}{\partial \theta} \right| &= \left| \int_{\mathbf{R}} \int_{\mu_2^*}^{\infty} \frac{\partial f(\theta | \mu_2)}{\partial \theta} f(\mu_2 | \mu_1) J(\mu_1) d\mu_{2,1} \right| \\ &\leq \int_{\mathbf{R}} \int_{\mu_2^*}^{\infty} \left| \frac{\partial f(\theta | \mu_2)}{\partial \theta} \right| f(\mu_2 | \mu_1) d\mu_{2,1} \\ &= \int_{\mu_2^*}^{\infty} \left| \frac{\partial f(\theta | \mu_2)}{\partial \theta} \right| d\mu_{2,1} \\ &= \int_{\mu_2^*}^{\infty} \left| -\frac{\partial f(\theta | \mu_2)}{\partial \mu_2} \right| d\mu_{2,1} \end{aligned}$$

²³See the beginning of Appendix B.1 for the definition of $p(\theta; \tilde{\theta}, J)$.

$$\leq 2 \max_{\mu_2 \in \mathbf{R}} f(\mu_2 \mid \theta) \leq \sqrt{\frac{4}{\pi}} \max\{\sqrt{\tau_1}, \sqrt{\tau}\}.$$

Therefore, we have

$$\frac{\partial p(\theta; \tilde{\theta}, J)}{\partial \theta} \geq -\left(\sqrt{\frac{2}{\pi}} + \sqrt{\frac{4}{\pi}}\right) \max\{\sqrt{\tau_1}, \sqrt{\tau}\} = -\sqrt{\frac{6+4\sqrt{2}}{\pi}} \max\{\sqrt{\tau_1}, \sqrt{\tau}\}.$$

Given the condition stated in Lemma A9, we have

$$\begin{aligned} \frac{\partial R(\theta, p(\theta; \tilde{\theta}, J))}{\partial \theta} &= R_\theta(\theta, p(\theta; \tilde{\theta}, J)) + R_p(\theta, p(\theta; \tilde{\theta}, J)) \frac{\partial p(\theta; \tilde{\theta}, J)}{\partial \theta} \\ &\geq \underline{R}_\theta - \bar{R}_p \cdot \sqrt{\frac{6+4\sqrt{2}}{\pi}} \max\{\sqrt{\tau_1}, \sqrt{\tau}\} > 0. \end{aligned}$$

Thus, $\text{sgn}[R(\theta, p(\theta; \tilde{\theta}, J))] = \text{sgn}(\theta - \tilde{\theta})$ for all $\theta \in \mathbf{R}$, which implies that $\tilde{\theta}$ is the regime-change threshold of a threshold equilibrium induced by J . $\square$

APPENDIX C: PROOFS OF RESULTS IN SECTION 5

C.1 Sketch of the proof of Theorem 2

The proof of Theorem 2 is a generalization of the proof of Theorem 1. In particular, we introduce a series of auxiliary games $\{\mathcal{G}_m\}_{m=1}^N$, where in game $\mathcal{G}_m$ the first m private signals are free to all agents. As in the proof of Theorem 1, $\mathcal{G}_1$ corresponds to the extended model as specified in Section 5 and $\mathcal{G}_N$ (where agents observe all signals for free) is equivalent to a static global game model in which agents observe a single signal with precision $N\tau$. Denote by $p(\hat{\theta}; \hat{\theta}, m)$ the aggregate attack in $\mathcal{G}_m$ when $\theta = \hat{\theta}$ and every agent takes $\hat{\theta}$ as the regime-change threshold.

Since $p(\hat{\theta}; \hat{\theta}, N) = 1 - \gamma$, it is enough to establish that $p(\hat{\theta}; \hat{\theta}, m) = p(\hat{\theta}; \hat{\theta}, m-1)$ for all $2 \leq m \leq N$. The desired result will then follow from backward induction on m . To establish that $p(\hat{\theta}; \hat{\theta}, m) = p(\hat{\theta}; \hat{\theta}, m-1)$, one can use essentially the same steps as those employed in the proof of Theorem 1. The key step is to establish a recursive form for the derivative of each of the value functions $T_m(\cdot)$ and $Q_m(\cdot)$, the counterparts of $T(\cdot)$ and $Q(\cdot)$ in $\mathcal{G}_m$. Such recursive forms are then used to show that the net effect of reducing a free signal (i.e., from $\mathcal{G}_m$ to $\mathcal{G}_{m-1}$) on the aggregate attack is 0. For more details, see the working paper version (<https://drive.google.com/file/d/1lXIRbB0eQILYZchfUQxeS-4avDiVeFf3/view?usp=sharing>) of this article.

C.2 Proof of Proposition 5

Part (i) The planner's solution to the N -signal case is in the same spirit as that of the 2-signal case; that is, the planner will offer an extra signal to an agent (if feasible) only if the agent will choose to not attack the regime without any further information. Based

on this characterization, we have²⁴

$$\begin{aligned}
 & 1 - \bar{p}(N) \\
 &= \int_{-\infty}^{\Phi^{-1}(\gamma)} \int_{-\infty}^{\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma)} \cdots \int_{-\infty}^{\sqrt{\frac{\tau_N}{\tau}} \Phi^{-1}(\gamma)} \left[\prod_{m=3}^N \phi(\nu_m - \nu_{m-1}) \right] \phi\left(\nu_2 - \sqrt{\frac{\tau_1}{\tau}} \nu_1\right) \\
 &\quad \times \phi(\nu_1) d\nu_{N \downarrow 1} \\
 &= \int_{-\infty}^{\Phi^{-1}(\gamma)} \int_{-\infty}^{\sqrt{\frac{\tau_2}{\tau}} \Phi^{-1}(\gamma)} \cdots \int_{-\infty}^{\sqrt{\frac{\tau_N}{\tau}} \Phi^{-1}(\gamma)} \int_{\mathbf{R}} \left[\prod_{m=3}^{N+1} \phi(\nu_m - \nu_{m-1}) \right] \phi\left(\nu_2 - \sqrt{\frac{\tau_1}{\tau}} \nu_1\right) \\
 &\quad \times \phi(\nu_1) d\nu_{N+1 \downarrow 1} \\
 &> \int_{-\infty}^{\Phi^{-1}(\gamma)} \cdots \int_{-\infty}^{\sqrt{\frac{\tau_N}{\tau}} \Phi^{-1}(\gamma)} \int_{-\infty}^{\sqrt{\frac{\tau_{N+1}}{\tau}} \Phi^{-1}(\gamma)} \left[\prod_{m=2}^{N+1} \phi(\nu_m - \nu_{m-1}) \right] \phi\left(\nu_2 - \sqrt{\frac{\tau_1}{\tau}} \nu_1\right) \\
 &\quad \times \phi(\nu_1) d\nu_{N+1 \downarrow 1} \\
 &= 1 - \bar{p}(N+1),
 \end{aligned}$$

which implies that $\bar{p}(N+1) > \bar{p}(N)$. Since $R(\underline{\theta}(N), \bar{p}(N)) = 0 = R(\underline{\theta}(N+1), \bar{p}(N+1))$, one concludes from the monotonicity of $R(\cdot, \cdot)$ that $\underline{\theta}(N+1) < \underline{\theta}(N)$.

Part (ii) When $N \rightarrow \infty$, using the same argument as in the proof of Lemma 1 in Dasgupta, Steiner, and Stewart (2012), one can show that for all $\hat{\theta} \in \mathbf{R}$, $\lim_{N \rightarrow \infty} \bar{p}(N) = 1$. Therefore, $\underline{\theta}^* \downarrow \underline{\theta}$ as $N \rightarrow \infty$.

REFERENCES

- Ahnert, Toni and Ali Kakhbod (2017), “Information choice and amplification of financial crises.” *The Review of Financial Studies*, 30 (6), 2130–2178. [0135]
- 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. [0135]
- Angeletos, George-Marios and Chen Lian (2016), “Incomplete information in macroeconomics: Accommodating frictions in coordination.” In *Handbook of Macroeconomics*, Vol. 2, 1065–1240, Elsevier. [0132, 0133]
- Angeletos, George-Marios and Alessandro Pavan (2007), “Efficient use of information and social value of information.” *Econometrica*, 75 (4), 1103–1142. [0135]
- Chernoff, Herman (1961), “Sequential tests for the mean of a normal distribution.” In *Proceedings of the Fourth Berkeley Symposium on Mathematical Statistics and Probability*, Vol. 1 (Jerzy Neyman, ed.), 79–91, University of California Press.

²⁴Here we write $d\nu_n d\nu_{n-1} \cdots d\nu_1$ as $d\nu_{n \downarrow 1}$ for notational simplicity. Also, $\tau_n = \tau_1 + (n-1)\tau$ for all $n \geq 2$.

- Colombo, Luca, Gianluca Femminis, and Alessandro Pavan (2014), "Information acquisition and welfare." *The Review of Economic Studies*, 81 (4), 1438–1483. [0132, 0133, 0135, 0143]
- Cooper, Russell (1999), *Coordination Games*. Cambridge University Press. [0132]
- Cooper, Russell and Andrew John (1988), "Coordinating coordination failures in Keynesian models." *The Quarterly Journal of Economics*, 103 (3), 441–463. [0132]
- Dasgupta, Amil (2007), "Coordination and delay in global games." *Journal of Economic Theory*, 134 (1), 195–225. [0134, 0135, 0144]
- Dasgupta, Amil, Jakub Steiner, and Colin Stewart (2012), "Dynamic coordination with individual learning." *Games and Economic Behavior*, 74 (1), 83–101. [0135, 0146, 0147, 0163]
- DeGroot, Morris H. (2005), *Optimal Statistical Decisions*. Wiley. [0135]
- Edmond, Chris (2013), "Information manipulation, coordination, and regime change." *Review of Economic Studies*, 80 (4), 1422–1458. [0136]
- Gittins, John C. (1979), "Bandit processes and dynamic allocation indices." *Journal of the Royal Statistical Society: Series B (Methodological)*, 41 (2), 148–164. [0135]
- Goldstein, Itay and Ady Pauzner (2005), "Demand-deposit contracts and the probability of bank runs." *The Journal of Finance*, LX (3), 1293–1327. [0136]
- Guimaraes, Bernardo and Stephen Morris (2007), "Risk and wealth in a model of self-fulfilling currency attacks." *Journal of Monetary Economics*, 54 (8), 2205–2230. [0135]
- Hellwig, Christian (2002), "Public information, private information, and the multiplicity of equilibria in coordination games." *Journal of Economic Theory*, 107 (2), 191–222. [0135]
- Hellwig, Christian and Laura Veldkamp (2009), "Knowing what others know: Coordination motives in information acquisition." *The Review of Economic Studies*, 76 (1), 223–251. [0132, 0135]
- Iachan, Felipe S. and Plamen T. Nenov (2015), "Information quality and crises in regime-change games." *Journal of Economic Theory*, 158, 739–768. [0135, 0148]
- Ke, T. Tony and J. Miguel Villas-Boas (2019), "Optimal learning before choice." *Journal of Economic Theory*, 180, 383–437. [0135]
- Liao, Xiaoye (2021), "Bayesian persuasion with optimal learning." *Journal of Mathematical Economics*, 97, 102534. [0136]
- Mathevet, Laurent and Jakub Steiner (2013), "Tractable dynamic global games and applications." *Journal of Economic Theory*, 148 (6), 2583–2619. [0135]
- McCall, John Joseph (1970), "Economics of information and job search." *The Quarterly Journal of Economics*, 84 (1), 113–126. [0135]

Morris, Stephen and Hyun Song Shin (1998), “Unique equilibrium in a model of self-fulfilling currency attacks.” *American Economic Review*, 88 (3), 587–597. [0132, 0135, 0136]

Morris, Stephen and Hyun Song Shin (2002), “Social value of public information.” *American Economic Review*, 92 (5), 1521–1534. [0132, 0135]

Morris, Stephen and Hyun Song Shin (2003), “Global games: Theory and applications.” In *Advances in Economics and Econometrics: Theory and Applications, Eighth World Congress* (M. Dewatripont, L. P. Hansen, and S. Turnovsky, eds.), 56–114, Cambridge University Press. [0133]

Morris, Stephen and Hyun Song Shin (2004), “Coordination risk and the price of debt.” *European Economic Review*, 48 (1), 133–153. [0135, 0159]

Moscarini, Giuseppe and Lones Smith (2001), “The optimal level of experimentation.” *Econometrica*, 69 (6), 1629–1644. [0135]

Myatt, David P. and Chris Wallace (2012), “Endogenous information acquisition in coordination games.” *The Review of Economic Studies*, 79 (1), 340–374. [0132, 0135]

Roberts, Kevin and Martin L. Weitzman (1981), “Funding criteria for research, development, and exploration projects.” *Econometrica*, 49 (5), 1261–1288. [0135]

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. [0136]

Rothschild, Michael (1974), “A two-armed bandit theory of market pricing.” *Journal of Economic Theory*, 9 (2), 185–202. [0135]

Sakovics, Jozsef and Jakub Steiner (2012), “Who matters in coordination problems?” *The American Economic Review*, 102 (7), 3439–3461. [0146]

Steiner, Jakub (2008), “Coordination cycles.” *Games and Economic Behavior*, 63 (1), 308–327. [0135]

Szkup, Michal (2022), “Preventing self-fulfilling debt crises: A global games approach.” *Review of Economic Dynamics*, 43, 22–55. [0136]

Szkup, Michal and Isabel Trevino (2015), “Information acquisition in global games of regime change.” *Journal of Economic Theory*, 160 (1), 387–428. [0133, 0135, 0144, 0148]

Ui, Takashi and Yasunori Yoshizawa (2015), “Characterizing social value of information.” *Journal of Economic Theory*, 158, 507–535. [0135]

Vives, Xavier (2014), “Strategic complementarity, fragility, and regulation.” *Review of Financial Studies*, 27 (12), 3547–3592. [0133, 0136, 0143, 0148]

Weitzman, Martin L. (1979), “Optimal search for the best alternative.” *Econometrica*, 47 (3), 641–654. [0135]

Yang, Ming (2015), “Coordination with flexible information acquisition.” *Journal of Economic Theory*, 158, 721–738. [0133, 0135]

Co-editor Bruno Strulovici handled this manuscript.

Manuscript received 7 November, 2023; final version accepted 1 July, 2025; available online 22 July, 2025.

All authors assume responsibility for all aspects of the paper.