

Diffusing Coordination Risk

Author(s): Deepal Basak and Zhen Zhou

Source: *The American Economic Review*, JANUARY 2020, Vol. 110, No. 1 (JANUARY 2020), pp. 271-297

Published by: American Economic Association

Stable URL: <https://www.jstor.org/stable/10.2307/26863280>

REFERENCES

Linked references are available on JSTOR for this article:

https://www.jstor.org/stable/10.2307/26863280?seq=1&cid=pdf-reference#references_tab_contents

You may need to log in to JSTOR to access the linked references.

JSTOR is a not-for-profit service that helps scholars, researchers, and students discover, use, and build upon a wide range of content in a trusted digital archive. We use information technology and tools to increase productivity and facilitate new forms of scholarship. For more information about JSTOR, please contact support@jstor.org.

Your use of the JSTOR archive indicates your acceptance of the Terms & Conditions of Use, available at <https://about.jstor.org/terms>



American Economic Association is collaborating with JSTOR to digitize, preserve and extend access to *The American Economic Review*

Diffusing Coordination Risk[†]

By DEEPAL BASAK AND ZHEN ZHOU*

In a regime change game, privately informed agents sequentially decide whether to attack without observing others' previous actions. To dissuade them from attacking, a principal adopts a dynamic information disclosure policy, frequent viability tests. A viability test publicly discloses whether the regime has survived the previous attacks. When such tests are sufficiently frequent, in the unique cutoff equilibrium, agents never attack if the regime passes the latest test, regardless of their private signals. We apply this theory to demonstrate that a borrower can eliminate panic-based runs by sufficiently diffusing the rollover choices across different maturity dates. (JEL C72, D82, G21)

In a coordination game, the strategic uncertainty that agents face concerning the actions and beliefs of others may lead to undesirable outcomes. Consider a borrower who has issued short-term debts to finance some illiquid investment. When the debt matures, a creditor may not roll over if he is wary of other creditors withdrawing their funds. This debt structure could cause a debt run solely on the basis of panic and not on underlying fundamental. We model this as a global game of regime change.¹ In a dynamic setting, we propose a simple information disclosure policy that influences the beliefs, and thus actions, of agents. This policy resolves strategic uncertainty and avoids undesirable outcomes.

Consider a regime, a principal, and a mass 1 of agents. The agents move sequentially: an agent $i \in [0, 1]$ moves at time i and decides whether to attack a regime, but he does not see other agents' previous actions. The underlying fundamental strength of the regime is θ . If θ is sufficient to withstand the aggregate attack, the

* Basak: Indian School of Business (email: deepal_basak@isb.edu); Zhou: PBC School of Finance, Tsinghua

¹ For examples of a global game of regime change, see Morris and Shin (1998) and Angeletos, Hellwig, and Pavan (2007) for the currency crisis, Goldstein and Pauzner (2005) for self-fulfilling bank runs, Vives (2014) for financial fragility, Edmond (2013) for riots and political change, and Konrad and Stolper (2016) for fight against tax havens. For more recent developments, see Szkup and Trevino (2015).

regime survives; otherwise, it does not. Suppose there is a threshold p such that an agent will not attack if and only if he believes with probability higher than p that the regime will survive. As is the standard in literature of global games, agents are uncertain about the underlying fundamental θ and receive some noisy private signals s_i .

The principal wants the regime to survive. However, she cannot influence the exogenous fundamental. She can strategically disclose information at different dates regarding the exogenous fundamental and past endogenous attacks to dissuade the agents from attacking.

First, let us consider two simple policies: no disclosure and full disclosure. Under no disclosure, agents have no information about past actions of other agents. Therefore, the game is essentially a simultaneous move game (see Morris and Shin 2003), and there is a unique equilibrium in which agents attack if and only if they receive a private signal below some cutoff.

On the other hand, if agents know the underlying fundamental and perfectly observe other agents' previous actions, all attacking is a possible equilibrium outcome. If there are only finitely many agents, it follows from backward induction that agents will coordinate on the payoff-dominant action.² However, in many examples such as the currency crisis problem (see Morris and Shin 1998 and Angeletos, Hellwig, and Pavan 2007), attacking the fixed exchange rate regime is the payoff-dominant action for agents, while the principal's interest is not aligned with that of the agents because the principal wants to defend it.

Thus, even if full disclosure is feasible, the principal may want to conceal some information. A vast range of partial disclosure policies exist. We consider a simple dynamic information disclosure policy called, *frequent viability tests*. As the name suggests, a viability test at some date $t \in [0, 1]$ checks whether the regime continues to be viable, i.e., if it has survived attacks thus far (if any). The result of this test is publicly disclosed. The principal only chooses an integer J denoting the frequency of viability tests. The tests are conducted at a regular interval of $1/J$, starting at 0.

Our main result shows that if the principal runs viability tests with sufficient frequency, then there is a unique cutoff equilibrium in which agents ignore their private information and never attack a viable regime.

In the spirit of Bayesian persuasion, we can interpret this policy as a recommendation by the principal to the agents not to attack a regime that passes the latest viability test and to attack if it fails. If the regime fails a viability test, it cannot survive even without any further attack. Hence, it is the dominant strategy for all agents in the subsequent groups to follow the principal's recommendation and attack. The challenging case arises when the regime passes the test.

If the regime passes a viability test, then the agents become more optimistic about the fundamental θ and other agents' beliefs about θ and so on. If one agent is less likely to attack, then it follows from strategic complementarity that others are also less likely to attack. Thus, it is intuitive that positive viability news reduces aggregate attack in the equilibrium.

²We discuss the case of full disclosure in detail in Section V.

When the principal runs viability tests for J times, the policy separates the agents into J groups. The $1/J$ mass of agents moving between the j th test and the $(j + 1)$ th test are referred to as the group j agents. We examine the equilibrium in cutoff strategies: for any j , after learning that the regime has passed the j th viability test, any agent in the group j does not attack if and only if his private signal is higher than some cutoff $\hat{s}_j$. Thus, when the regime is stronger (higher θ), more agents receive signals that exceed the cutoff and fewer agents from each group attack. This induces a nondecreasing sequence of fundamental cutoffs $\{\underline{\theta}_{j-1}\}_{j=1}^J$ such that the regime passes the j th test if and only if the underlying fundamental is no lower than the cutoff $\underline{\theta}_{j-1}$.

These cutoffs are endogenous, which makes the policy history-dependent. Moreover, the effectiveness of the present viability news depends on the effectiveness of the future viability news. We develop an inductive argument: if the tests are sufficiently frequent, then agents in the group j follow the principal's recommendation if they believe that agents in the subsequent groups will do so regardless of the cutoff strategy played by others in the past (equivalently for any $\underline{\theta}_{j-1}$).

Given that the regime has passed the j th viability test and the agents in subsequent groups will follow the principal's recommendation, the regime will survive as long as it withstands attacks from the group j . Suppose that the agents in group j follow some cutoff strategy $\hat{s}_j$. This monotone strategy can constitute an equilibrium if the "marginal agent" with the signal $\hat{s}_j$ is indifferent between attacking and not attacking, i.e., he believes that the regime survives with probability p . We argue that there exists a $\hat{J}$ such that if the group size is smaller than $1/\hat{J}$, then no such equilibrium can exist.

This is because, under frequent viability tests, the size of each group is small. If the group size is smaller, the magnitude of attack from this group is less for any given cutoff strategy $\hat{s}_j$. Given that no one from the subsequent groups will attack a viable regime, the regime is more likely to survive. This, in turn, amplifies the influence of positive news from the j th viability test and makes the marginal agent more confident about the survival of the regime. We demonstrate that the marginal agent's belief regarding the chance of survival *uniformly converges* to 1 with group size regardless of $\underline{\theta}_{j-1}$. It follows from uniform convergence that there exists a $\hat{J}$ such that, under sufficiently frequent viability tests (i.e., $J > \hat{J}$), the marginal agent always believes that the regime will survive with a probability strictly higher than p ; and thus, he strictly prefers not to attack if agents in the subsequent groups will not attack a regime that passes the latest viability test.

Under sufficiently frequent viability tests, since no one moves after the last group of agents, there cannot be any cutoff equilibrium in which group J agents attack a viable regime (regardless of the cutoff strategies played by others in the past)

sound obvious and not a deliberate attempt to manipulate agents' beliefs. One may believe that the news is not likely to have any substantial influence on creditors' behavior. However, the borrower provides this news J many times when only $1/J$ fraction of debts mature at a time. We extend our model to show that a sufficient asynchronous debt structure (sufficiently large J) makes the borrower immune to panic-based runs.

Related Literature.—Similar to a Pigouvian planner, our principal attempts to achieve the desired outcome that the market fails to deliver. Sákovics and Steiner (2012) and Cong, Grenadier, and Hu (forthcoming) find optimal subsidies that, at a given cost, maximize the likelihood of successful coordination.³ Unlike the planner in the papers mentioned above, we consider a principal who cannot offer monetary incentives. Instead, she acts as an information designer and discloses some relevant information. We consider a canonical global game of regime change, and as is standard in the literature, we assume a private information environment (see Carlsson and van Damme 1993). Similar to Bergemann and Morris (2013), the principal does not have access to agents' private information.

The two most closely related papers are Inostroza and Pavan (2018) and Goldstein and Huang (2016). The authors also propose a partial information disclosure policy, a one-time “stress test.”⁴ In this paper, the principal exploits the fact that agents do not move simultaneously and runs multiple viability tests over time. Thus, the paper belongs to the recent literature on dynamic information design. Ely (2017) is the first paper to extend the static Bayesian persuasion model of Kamenica and Gentzkow (2011) to the dynamic setting, where the disclosure policies are history-independent. Doval and Ely (2019) considers a general extensive form game with incomplete information in which the principal does not know the exact extensive form that governs the play.⁵

Our paper is related to the dynamic coordination game literature as well. Dasgupta (2007) considers a two-period problem in which agents receive noisy private information about the attacks from the first period. Similar to our model, in Angeletos, Hellwig, and Pavan (2007) and Huang (2017) agents learn from the past through viability news, which results in multiple cutoff equilibria. Unlike the papers mentioned above, we optimally choose the information transmission over time. Frankel and Pauzner (2000) introduces asynchronicity by allowing the agents to revise their decisions following a Poisson process. In their model, agents know the current state but are uncertain about the volatile future.⁶ He and Xiong (2012) extends this framework to study the role of volatile fundamental under asynchronous debt structure.

³ While Sákovics and Steiner (2012) considers heterogeneous agents and demonstrates that subsidizing the more reluctant agents matters more, Cong, Grenadier, and Hu (forthcoming) demonstrates that if liquidity injection is equally costly across periods in a dynamic coordination problem, early injection is more helpful.

⁴ Information design has been studied in other strategic contexts, such as voting and auctions. See Bergemann and Morris (2019) for a recent survey of this literature.

⁵ Significant work has been conducted on dynamic information feedback in the context of strategic experimentation. See Hörner and Skrzypacz (2016) for a survey of this literature.

⁶ Sequential move has also been studied in the context where agents' payoffs are independent of others' actions but there is incomplete information about a common fundamental (see Banerjee 1992 and Bikhchandani, Hirshleifer, and Welch 1992). For games of strategic substitutes such as public good contribution, Varian (1994) shows that sequential move usually reduces the total contribution.

While the two papers mentioned above take the asynchronous structure as a given, we take a step back and address what necessitates the asynchronous debt structure in the first place. We find the justification in terms of information transmission and identify the optimal design of the asynchronous structure from the borrower's perspective.

The rest of the paper is organized as follows. Section I describes the model, the diffused policy, and the solution concept. Section II illustrates the role of a one-time viability test. In Section III, we leverage this idea to develop the argument for frequent viability tests and establish our main result. In Section IV, we extend the model to study its application in debt rollover. In Section V, we discuss full disclosure policy and the factors that may make persuasion harder or even impossible. Section VI concludes this paper. Proofs that are not presented in the paper are provided in the Appendix.

I. A Simple Model of Regime Change

There is a regime, a principal, and a mass 1 of agents. The agents are indexed by $i \in [0, 1]$. Agent i moves at time i and takes an action $a_i \in \{0, 1\}$, where 1 (0) indicates attacking (not attacking) the regime.

The agents cannot observe other agents' previous actions. Let θ be the underlying fundamental strength of the regime and $\mathfrak{w} = \int_0^1 a_i di$ be the aggregate attack. The regime survives if and only if its fundamental strength θ is strong enough to withstand the aggregate attack against it, i.e., $\theta \geq \mathfrak{w}$.

Agents are assumed to be ex ante identical and risk neutral. Given that the fundamental strength (θ) and the aggregate attack ($\mathfrak{w}$), let $u(a_i, \mathfrak{w}, \theta)$ be the payoff for agent i if he takes action a_i , where

$$(1) \quad u(0, \mathfrak{w}, \theta) = \begin{cases} b_1 & \text{if } \theta \geq \mathfrak{w} \\ c_0 & \text{if } \theta < \mathfrak{w} \end{cases}, \quad u(1, \mathfrak{w}, \theta) = \begin{cases} c_1 & \text{if } \theta \geq \mathfrak{w} \\ b_0 & \text{if } \theta < \mathfrak{w} \end{cases}.$$

We assume that (i) $b_1 > c_1$, i.e., if an agent knows that the regime will survive, not attacking is the desirable action; and (ii) $b_0 > c_0$, i.e., if an agent knows that the regime will not survive, attacking is the desirable action.⁷

However, agents are uncertain about θ and hence do not know whether the regime will survive. Each agent i receives a noisy private signal about θ , denoted by $s_i = \theta + \sigma \epsilon_i$, where ϵ_i is a random noise with zero mean, and $\sigma >$

$\bar{\theta} > 1 + \sigma$. Distribution Π admits a continuous density π that is strictly positive and bounded in $[-\sigma, 1 + \sigma]$, i.e., $0 < \underline{\pi} \leq \pi(\theta) \leq \bar{\pi} < \infty$.

If $\theta \geq 1$, the regime will surely survive regardless of w , and if $\theta < 0$, the regime cannot survive regardless of w . Only when $\theta \in [0, 1]$, the regime's survival is dependent upon the aggregate attack w . Note that for any $\theta \in [0, 1]$, the proportion of agents who receive a signal lower than s is $F(\sqrt{\tau}(s - \theta))$. For an agent with private signal $s \geq 1 + \sigma/2$, the dominant strategy is not to attack, and for an agent with private signal $s < -\sigma/2$, the dominant strategy is to attack. An agent who receives a private signal $s \in [-\sigma/2, 1 + \sigma/2]$, does not have a dominant strategy. He updates his belief that $\theta \geq A$ (for some A) as⁹

$$\Pr(\theta \geq A|s) = \int_A^{\bar{\theta}} \left(\frac{\pi(\theta)f(\sqrt{\tau}(s - \theta))}{\int_{\underline{\theta}}^{\bar{\theta}} \pi(\theta)f(\sqrt{\tau}(s - \theta))d\theta} \right) d\theta.$$

The log-concavity of f is equivalent to the *monotone likelihood ratio property* (MLRP). That is, for $s_1 > s_2$,

$$\frac{f(\sqrt{\tau}(s_1 - \theta))}{f(\sqrt{\tau}(s_2 - \theta))}$$

is increasing in θ . This implies that, for any $A > B$, $\Pr(\theta \geq A|s, \theta \geq B)$ is increasing in s (see the online Appendix for the formal argument).

Suppose an agent believes that the regime survives with probability P . The expected payoff from attacking ($a_i = 1$) is $Pc_1 + (1 - P)b_0$, while the expected payoff from not attacking ($a_i = 0$) is $Pb_1 + (1 - P)c_0$. Therefore, the agent does not attack if and only if he believes that the probability of the regime surviving (or P) is greater than

$$(2) \quad p := \frac{1}{1 + \frac{b_1 - c_1}{b_0 - c_0}}.$$

The Principal's Objective.—The principal obtains a payoff of 1 if the regime survives and 0 otherwise. If $\theta < 0$, the regime cannot survive. Otherwise, the regime is viable, which means the regime survives if no one attacks it. However, a viable regime may not survive if agents attack. We call the ex ante probability that a viable regime may not survive: *coordination risk*. The principal seeks to minimize this coordination risk. If agents never attack a viable regime, then this risk is eliminated. However, an agent will not attack only if he believes that the regime will survive with a probability of at least p . We denote p the reluctance of the agents. Note that when attacking is the payoff-dominant action ($b_1 < b_0$), the principal's interest is not aligned with that of the agents (and the agents are more reluctant). We discuss this in Section V.

⁹Note that although we use the $\underline{\theta}$ and $\bar{\theta}$ in the limit of the integration, for a given s , if $\theta > s + \sigma/2$ or $\theta < s - \sigma/2$, $f(\sqrt{\tau}(s - \theta)) = 0$.

A Diffused Policy J.—The principal, who can be thought of as an information designer, can disclose some information over time based on the (exogenous) fundamental as well as the (endogenous) history of actions. We focus on a particular dynamic information disclosure policy. Consider a test that checks whether the regime continues to be viable at any date t , i.e., if it can survive in the absence of further attacks, and disclose the test result publicly. We call such tests—viability tests, and a positive test result—the public news of continued viability (PNV). Before agents move and the fundamental θ is realized, the principal chooses the frequency at which such viability tests should run. The policy is announced publicly. Without loss of generality, we only consider regular intervals between tests.

DEFINITION 1: *A diffused policy J conducts the viability test at a regular interval of $\alpha = 1/J$ starting at 0, i.e., at $(j - 1)\alpha$ for $j = 1, 2, \dots, J$.*

The j th test is conducted at time $(j - 1)\alpha$ for any $j = 1, 2, \dots, J$. The agents in $[(j - 1)\alpha, j\alpha)$ move between the j th test and the next test. We refer to these α mass of agents as group j . The policy J separates the agents into J different groups based on the latest public news from the viability tests they can receive. A more diffused policy indicates the increased frequency of viability tests (higher J), or equivalently the group size α is smaller.

Cutoff Equilibrium.—Under a diffused policy J , the group j agents obtain public information on whether the regime has passed the j th viability test before taking action. Furthermore, each agent has a private signal regarding the underlying fundamental. If the regime fails a viability test, then the regime will not survive regardless of the remaining agents' actions. Hence, attacking is the dominant strategy for the group j agents. The strategic decision is nontrivial only when the regime passes the test, which is the focus of our analysis. In what follows, we limit our attention to the perfect Bayesian equilibrium in monotone strategies such as the case after agents learn that the regime has passed the viability test, the probability that an agent i attacks $\rho_i(s_i)$ is nonincreasing in the private signal s_i . When the agents follow monotone strategies, a larger mass of agents attack against a weaker regime. Therefore, a stronger regime is more likely to survive. Under log-concavity of f , agent i believes the regime is strictly more likely to survive when s_i is higher. Thus, in equilibrium, all agents in the same group j (for any j) must follow a symmetric cutoff strategy: attack if and only if $s_i < \hat{s}_j$.¹⁰ We refer to such equilibrium as cutoff equilibrium.

As the game continues, more agents attack. Hence, there exists a nondecreasing sequence of cutoff $\{\underline{\theta}_{j-1}\}_{j=1}^J$ such that the regime passes the j th viability tests if and only if $\theta \geq \underline{\theta}_{j-1}$. Let us define $\hat{\theta}$ such that if $\theta \geq \hat{\theta}$, then the regime passes all viability tests and survives. If there were a $(J + 1)$ th viability test after group J , then $\underline{\theta}_J$ would be equal to $\hat{\theta}$.

¹⁰ Suppose an agent in group j randomizes upon receiving a signal $\hat{s}_j$. Then, the agent is indifferent between attacking and not attacking. Therefore, any agent i in group j must play $\rho_i(s_i) = 0$ for $s_i > \hat{s}_j$ and $\rho_i(s_i) = 1$ for s_i

Persuasion.—In the spirit of Bayesian persuasion, one can interpret this policy as a recommendation by the principal to the agents to “not attack” when the regime passes the test and “attack” when it does not.

DEFINITION 2: *A policy is persuasive if there is no cutoff equilibrium, in which an agent attacks a regime that passes the latest test.*

If a diffused policy is persuasive, then the only cutoff equilibrium is that in which agents ignore their private information and follow the principal’s recommendation. This means that no agent attacks a viable regime and thus, any regime with underlying fundamental $\theta \geq 0$ survives. In other words, under a persuasive policy, the only equilibrium cutoff fundamental is $\hat{\theta} = 0$, and the coordination risk is eliminated.

II. One-Time Viability Test

Consider a degenerate policy, in which no tests are conducted. Since agents do not have any information about past actions of other agents, the game is essentially a simultaneous move regime change game. It follows from Morris and Shin (2003) that if the private signals are sufficiently precise, then there is a unique cutoff equilibrium. Under the uniform prior, one can explicitly solve for the unique cutoff signal as $s^* = p + (1/\sqrt{\tau})F^{-1}(p)$. This means that the regime survives the attacks if and only if $\theta \geq p$. Recall that an agent can be dissuaded from attacking if he believes the regime will survive with probability higher than p . Thus, it is intuitive that when p is higher, the ex ante chance of survival is lower.

To understand the effects of the viability test, let us first consider the policy $J = 1$, i.e., a one-time viability test. The regime passes the test when $\theta \geq 0$. The agents have different beliefs over θ based on their private signals. Thus, PNV will not affect all agents in the same fashion. The agents with private signal $s_i \in [-\sigma/2, \sigma/2)$ once believed that θ could be less than 0, but the public news tells them otherwise. This makes them more optimistic about the strength of the regime, and thus less likely to attack. Owing to the strategic complementarity, other agents, including those who already know that $\theta \geq 0$ from their private information, become more optimistic about the success of the regime; hence, they are also less likely to attack.¹¹ It is intuitive that after learning that the regime is viable, in equilibrium, the agents will attack less aggressively. However, the news may not be strong enough, and there can be an equilibrium in which an agent attacks the regime that has passed the viability test (as in Goldstein and Huang 2016).

Consider a hypothetical game that exactly reflects that presented above with one exception: it is played between some $\alpha < 1$ mass of agents instead of 1 mass of agents. The smaller the α , the more likely it is that $\theta \geq \alpha$; in other words, even if all the α mass of agents attack, the regime will survive. However, agents may receive private signal $s_i < \alpha - \sigma/2$ and believe that the regime will not survive if others attack. To understand whether such agents can be persuaded not to attack, we need to look into their beliefs about others in the equilibrium. Below, we investigate

¹¹ Note that PNV makes the agents more optimistic about the success of the regime regardless of whether attacking is the payoff-dominated action. We revisit this issue later in Section V.

how the equilibrium strategy is affected when a small mass of agents learns that the regime has passed the viability test.

Persuasive Viability Test

Although the positive news $\theta \geq 0$ may not be strong enough to persuade a mass 1 of agents, the following lemma shows that a one-time viability test can be persuasive when the mass of agents is sufficiently small. We emphasize this result because it plays a crucial role in deriving our main result.

LEMMA 1: *Suppose that there is only one group with size α and they learn PNV. There exists $\alpha^* > 0$ such that there is no cutoff equilibrium in which an agent attacks a viable regime if and only if $\alpha < \alpha^*$.*

PROOF:

Suppose that agents follow some cutoff strategy, $\hat{s}$. Then, for any θ , the aggregate attack is $\alpha \Pr(s < \hat{s} | \theta) = \alpha F(\sqrt{\tau}(\hat{s} - \theta))$. Define $A(\hat{s}, \alpha) \in [0, \alpha]$ such that

$$(A^\alpha) \quad \alpha F(\sqrt{\tau}(\hat{s} - A(\hat{s}, \alpha))) = A(\hat{s}, \alpha).$$

By definition, the regime survives if and only if $\theta \geq A(\hat{s}, \alpha)$. We refer to this as the aggregate condition (A^α) .

After learning PNV ($\theta \geq 0$), an agent with private signal s_i believes that the regime will survive with probability

$$(3) \quad \Pr(\theta \geq A(\hat{s}, \alpha) | s_i, \theta \geq 0) = \frac{\int_{A(\hat{s}, \alpha)}^{\bar{\theta}} \pi(\theta) f(\sqrt{\tau}(s_i - \theta)) d\theta}{\int_0^{\bar{\theta}} \pi(\theta) f(\sqrt{\tau}(s_i - \theta)) d\theta}.$$

It follows from log-concavity of f that this probability is higher for the agents with higher private information s_i (see the online Appendix for the formal proof). Define $I^v(\hat{\theta})$ such that if the survival criteria is $\theta \geq \hat{\theta}$, the agent who receives private signal $s_i = I^v(\hat{\theta})$ is indifferent between attacking and not attacking:

$$(I^v) \quad \frac{\int_{\hat{\theta}}^{\bar{\theta}} \pi(\theta) f(\sqrt{\tau}(I^v(\hat{\theta}) - \theta)) d\theta}{\int_0^{\bar{\theta}} \pi(\theta) f(\sqrt{\tau}(I^v(\hat{\theta}) - \theta)) d\theta} = p.$$

We refer to this as the indifference condition (I^v) .

Suppose that all agents are following the cutoff strategy $\hat{s}$ and consequently, the survival criterion is $\hat{\theta} = A(\hat{s}, \alpha)$. By definition, any agent with $s_i = \beta(\hat{s}) := I^v(A(\hat{s}, \alpha))$ is indifferent between attacking and not attacking. Hence, a cutoff strategy $\hat{s}$ can constitute an equilibrium if $\hat{s} = \beta(\hat{s})$. If that is the case, we have

$$\frac{\int_{A(\hat{s}, \alpha)}^{\bar{\theta}} \pi(\theta) f(\sqrt{\tau}(\hat{s} - \theta)) d\theta}{\int_0^{\bar{\theta}} \pi(\theta) f(\sqrt{\tau}(\hat{s} - \theta)) d\theta} = p.$$

Finally, substituting $\hat{s}$ from the aggregate condition (A^α) in the above, we get

$$\frac{\int_{A(\hat{s}, \alpha)}^{\bar{\theta}} \pi(\theta) f\left(\sqrt{\tau}(A(\hat{s}, \alpha) - \theta) + F^{-1}\left(\frac{A(\hat{s}, \alpha)}{\alpha}\right)\right) d\theta}{\int_0^{\bar{\theta}} \pi(\theta) f\left(\sqrt{\tau}(A(\hat{s}, \alpha) - \theta) + F^{-1}\left(\frac{A(\hat{s}, \alpha)}{\alpha}\right)\right) d\theta} = p.$$

The left-hand side captures the belief of the marginal agent, given he learns PNV, that “given that α fraction of agents play the cutoff strategy $\hat{s}$, the regime has a per capita fundamental strength that surpasses the required cutoff $A(\hat{s}, \alpha)/\alpha$, and thus survives.” We define $G: [0, 1] \times [0, 1] \rightarrow [0, 1]$ as follows:

$$(G) \quad G(x, \alpha) := \frac{\int_{\alpha x}^{\bar{\theta}} \pi(\theta) f\left(\sqrt{\tau}(\alpha x - \theta) + F^{-1}(x)\right) d\theta}{\int_0^{\bar{\theta}} \pi(\theta) f\left(\sqrt{\tau}(\alpha x - \theta) + F^{-1}(x)\right) d\theta}.$$

Note that $G(x = A(\hat{s}, \alpha)/\alpha, \alpha)$ represents the belief of the marginal agent. By definition, $G(x, \alpha = 0) = 1$ for any $x \in [0, 1]$, and $G(x = 0, \alpha) = 1$ for any $\alpha \in [0, 1]$. For any

]

The News Effect.—As we have already highlighted, PNV makes agents more optimistic about the survival of the regime. For any survival criteria $\hat{\theta}$, the private signal realization $I^v(\hat{\theta})$ that makes agents indifferent between attacking and not attacking, is lower than that in the absence of any viability test. To observe this, consider the signal $I(\hat{\theta})$ that makes an agent indifferent between attacking and not attacking when there is no viability test; $I(\hat{\theta})$ must satisfy

$$(I) \quad \frac{\int_{\hat{\theta}}^{\bar{\theta}} \pi(\theta) f(\sqrt{\tau}(I(\hat{\theta}) - \theta)) d\theta}{\int_{\underline{\theta}}^{\bar{\theta}} \pi(\theta) f(\sqrt{\tau}(I(\hat{\theta}) - \theta)) d\theta} = p.$$

It follows from log-concavity of f that the LHS of condition (I) (or (I^v)) is increasing in the required cutoff signal $I(\hat{\theta})$ (or $I^v(\hat{\theta})$). Compared with condition (I^v) , the denominator on the LHS of condition (I) is larger. Thus, it is easy to see that the required cutoff signal $I^v(\hat{\theta}) \leq I(\hat{\theta})$. By the same logic, one can easily check that the required cutoff signal $I^v(\hat{\theta})$ and $I(\hat{\theta})$ increase with the required fundamental strength $\hat{\theta}$.

The Group Size Effect.—It directly follows from the aggregate condition (A^α) that a small group size α would make the success criterion $A(\hat{s}, \alpha)$ lower, given any $\hat{s}$.

The Combined Effect.—Combining the two forces, we can say that for any $\hat{s}$, a smaller group size α translates into a lower $A(\hat{s}, \alpha)$ and in that case, PNV lowers $\beta(\hat{s}) = I^v(A(\hat{s}, \alpha))$, which is the cutoff signal that makes an agent indifferent between attacking and not attacking. If the group size α is sufficiently small, this combined effect is so significant that for any $\hat{s}$, $\beta(\hat{s}) < \hat{s}$.

Figure 1 explains this combined effect graphically. We use the uniform prior and triangle probability density of error: $\hat{f}(x) := (2 + 4x)\mathbf{1}(-0.5 \leq x < 0) + (2 - 4x)\mathbf{1}(0 \leq x \leq 0.5)$.¹² Consider any cutoff per capita fundamental x . As defined in equation (G), $G(x, \alpha)$ is the belief of the marginal agent that the regime will survive. Figure 1 plots this belief against any candidate cutoff $x \in [0, 1]$. Under a degenerate policy $G(x, \alpha)$ is the 45° line and hence at the cutoff per capita fundamental $x = p$, the agent with the cutoff signal is indifferent. Thus, under a uniform prior, in the absence of PNV, a small group size does not affect the equilibrium per capita fundamental cutoff.

We can see that PNV pushes this belief upward. As α decreases, $G(x, \alpha)$ increases. However, under the general prior, this monotonicity may not hold. Nevertheless, since $G(x, \alpha)$ uniformly converges to 1, there is an α^* such that for $\alpha < \alpha^*$, $G(x, \alpha) > p$ for all x . This implies that the only possible cutoff solution is $x = 0$ and the implied unique equilibrium is that in which no one attacks a viable regime.

A Lower Bound.—In the following section, we consider frequent viability tests that induce a dynamic setting. Constructing the belief of the marginal agent in this dynamic setting will be more involved. Below, we construct a lower bound $\underline{G}(x, \alpha)$

¹²Note that the error distribution

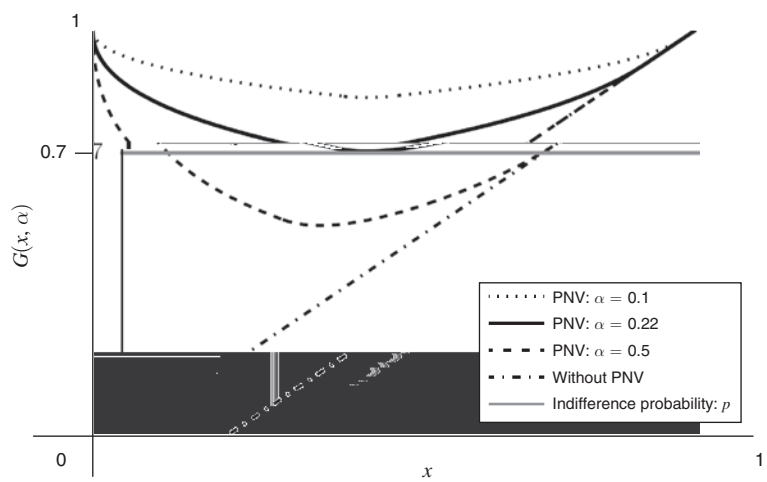


FIGURE 1. $G(x, \alpha)$ WHEN $\{\underline{\theta}, \bar{\theta}\} = \{-1, 2\}$, $\tau = 1$, $\pi(\theta)$ IS UNIFORM AND $f = \bar{f}$

on the belief $G(x, \alpha)$ of the marginal agent. This lower bound will be useful in the dynamic setting. Let us define

($\underline{G}$)
$$\underline{G}(x, \alpha) := 1 - \left(\frac{\bar{\pi} \bar{f}}{\underline{\pi} \underline{f}} \right) \left(\frac{\alpha x}{\alpha x + \frac{1}{\sqrt{\tau}} \left(F^{-1}(x) + \frac{1}{2} \right)} \right)$$

for $x > 0$ and $\underline{G}(0, \alpha) = 1$. It is clear that for any $x > 0$,¹³

$$G(x, \alpha) \geq 1 - \frac{\int_0^{\alpha x} \bar{\pi} \bar{f} d\theta}{\int_0^{\alpha x + \frac{1}{\sqrt{\tau}} F^{-1}(x) + \frac{1}{2\sqrt{\tau}}} \underline{\pi} \underline{f} d\theta} = \underline{G}(x, \alpha),$$

and $G(x = 0, \alpha) = \underline{G}(x = 0, \alpha) = 1$.

COROLLARY 1: *There exists $\hat{\alpha} \in (0, \alpha^*]$ such that whenever $\alpha < \hat{\alpha}$, $G(x, \alpha) \geq \underline{G}(x, \alpha) > p$ for all $x \in [0, 1]$.*

The Corollary 1 directly follows from the same argument as in Lemma 1 since $\underline{G}(x, \alpha)$ also uniformly converges to 1 as $\alpha \rightarrow 0$. Note, $\hat{\alpha} \leq \alpha^*$ because $\alpha < \alpha^*$ serves as the necessary and sufficient condition for Lemma 1 while $\alpha < \hat{\alpha}$ is only the sufficient condition. In other words, $\hat{\alpha}$ qualifies as a group size that is small enough to be dissuaded from attacking the regime by a viability test.

III. Main Result: Frequent Viability Tests

We understand that positive viability news can dissuade a sufficiently small mass of agents from attacking. Now, let us return to the frequent viability tests: a

¹³Note that, for any cutoff strategy $\hat{s} \in \left[-\frac{1}{2\sqrt{\tau}}, 1 + \frac{1}{2\sqrt{\tau}}\right)$, the effective upper bound for the integral is $\hat{s} + \frac{1}{2\sqrt{\tau}}$ instead of θ .

TABLE 1—CUTOFF EQUILIBRIA FOR $J = 2$

Equilibrium	$\hat{s}_1$	$\hat{s}_2$	$\underline{\theta}_1$	$\hat{\theta}$
1	−0.50	−0.50	0	0
2	−0.26	−0.46	0.04	0.04
3	0.46	−0.15	0.35	0.35
4	0.77	0.72	0.46	0.66
5	0.57	0.29	0.39	0.46

Note: Parameter values: $\tau = 1$, $p = 0.7$, π is uniform in $[-1, 2]$ and $f = \hat{f}$.

diffused policy J . There can be multiple cutoff equilibria. We relegate the full characterization of cutoff equilibria for any diffused policy J to the online Appendix. The following numerical example shows all possible cutoff equilibria for $J = 2$.

NUMERICAL EXAMPLE 1: *In Table 1, the first equilibrium is the one in which no agent attacks a viable regime ($\hat{s}_1 = \hat{s}_2 = -\sigma/2$). Thus, $\hat{\theta} = 0$. Similar to the one-time viability test, this remains a possible equilibrium outcome. There are two equilibria (2 and 3), in which no agent in group 2 attacks a regime if the regime passes the second viability test ($\hat{s}_2 = \underline{\theta}_1 - (\sigma/2)$). Thus, if the regime passes the second viability test, it survives ($\underline{\theta}_1 = \hat{\theta}$). It is also possible (equilibrium 4 and 5) that agents in both group 1 and 2 attack a viable regime.*

It is easy to see that for any diffused policy J , one can always construct a cutoff equilibrium in which for some $j = 1, 2, \dots, J$, no agent attacks the regime after the regime passes the j th viability test. Thus, the regime that passes the j th viability test continues to be viable until the end. However, if an agent believes that others may attack a regime even after it passes the viability tests, then he may do the same: particularly if this agent has received a very low private signal; hence, multiple cutoff equilibria may arise. Nevertheless, we argue that when the principal adopts a sufficiently diffused policy, the only possible cutoff equilibrium is the one in which all agents follow the principal’s recommendation and do not attack a viable regime.

THEOREM 1: *A diffused policy J with $J > \hat{J} = 1/\hat{\alpha}$ is persuasive.*

Consider the case in which group j agents learn that the regime has passed the j th viability test and decide whether to attack the regime. When $J > \hat{J}$, there are less than $\hat{\alpha}$ mass of agents in group j . Recall from Corollary 1 that a viability test can be persuasive if the mass of agents is smaller than $\hat{\alpha}$. However, there are two crucial differences in this dynamic case. First, whether the regime can pass the j th viability test depends on θ as well as past attacks. Thus, the interpretation of the positive viability news is history-dependent. To observe this, note that the viability test discloses whether $\theta \geq \underline{\theta}_{j-1}$, where $\underline{\theta}_{j-1}$ is endogenous. Second, although only α mass of agents are moving, $(1 - j\alpha)$ mass of agents will move later, and the regime may fail if agents in the subsequent groups attack. Thus, the effectiveness of j th viability test is dependent upon the effectiveness of future viability tests.

However, we build on Lemma 1 and Corollary 1, and argue that the following statement is true for any j , in particular for $j = 0$.

M_j : When the agents in any group $j' > j$ learn that the regime has passed the j' th viability test (regardless of whichever cutoff strategy was played by others in the past), any cutoff strategy in which an agent in group j' may attack violates sequential rationality.

As we only consider cutoff strategies, we sometimes refer to the statement above as “no agent in group $j' > j$ attacks a regime that passes the j' th viability test.” The theorem claims that when $J > \hat{J}$, $\mathbf{M}_0$ is true. We prove this using the induction argument; regardless of history, the agents in any group j follow the principal’s recommendation, if they believe that the agents in the subsequent groups will do the same. Formally speaking,

N_j : If M_j is true, then M_{j-1} is true.

Step A:

LEMMA 2: If $\alpha < \hat{\alpha}$

For any cutoff strategies the agents may have played in the past, PNV means $\theta \geq \underline{\theta}_{j-1}$, where $\underline{\theta}_{j-1}$ solves condition (4) with equality.¹⁴

Suppose that the agents in group j , follow a cutoff strategy $\hat{s}_j$. Let us define $A_j(\hat{s}_j, \alpha, \underline{\theta}_{j-1})$ such that

$$(A_j^\alpha) \quad A_j(\hat{s}_j, \alpha, \underline{\theta}_{j-1}) = \underline{\theta}_{j-1} + \alpha F\left(\sqrt{\tau}(\hat{s}_j - A_j(\hat{s}_j, \alpha, \underline{\theta}_{j-1}))\right).$$

Unless $\theta \geq \underline{\theta}_{j-1}$, the regime will not pass the j th viability test. Given the cutoff strategy $\hat{s}_j$, the second term on the RHS captures the aggregate attack from group j when $\theta = A_j$. Based on the condition (A_j^α) , if $\theta \geq A_j$, the regime will pass the j th viability test and then sustain the attack from group j .¹⁵

Given M_j is true, this means whenever $\theta \geq A_j$, the regime survives.

The marginal agent in group j with signal $\hat{s}_j$ believes that the regime will survive with probability at least

$$\Pr(\theta \geq A_j | \hat{s}_j, \theta \geq \underline{\theta}_{j-1}) = \frac{\int_{A_j}^{\bar{\theta}} \pi(\theta) f(\sqrt{\tau}(\hat{s}_j - \theta)) d\theta}{\int_{\underline{\theta}_{j-1}}^{\bar{\theta}} \pi(\theta) f(\sqrt{\tau}(\hat{s}_j - \theta)) d\theta}.$$

Substituting the aggregate condition (A_j^α) , we can write the marginal agent's belief as

$$\frac{\int_{A_j}^{\bar{\theta}} \pi(\theta) f\left(\sqrt{\tau}\left(A_j - \theta\right) + F^{-1}\left(\frac{A_j - \underline{\theta}_{j-1}}{\alpha}\right)\right) d\theta}{\int_{\underline{\theta}_{j-1}}^{\bar{\theta}} \pi(\theta) f\left(\sqrt{\tau}\left(A_j - \theta\right) + F^{-1}\left(\frac{A_j - \underline{\theta}_{j-1}}{\alpha}\right)\right) d\theta}.$$

Substituting the per capita required cutoff strength for agents in group j to pass the next test as $x = (A_j - \underline{\theta}_{j-1})/\alpha$, we get that belief of the marginal agent with signal $\hat{s}_j = (1/\sqrt{\tau})F^{-1}(x) + \alpha x + \underline{\theta}_{j-1}$ is

$$(5) \quad \frac{\int_{\underline{\theta}_{j-1} + \alpha x}^{\bar{\theta}} \pi(\theta) f(\sqrt{\tau}(\alpha x - \theta) + F^{-1}(x) + \sqrt{\tau}\underline{\theta}_{j-1}) d\theta}{\int_{\underline{\theta}_{j-1}}^{\bar{\theta}} \pi(\theta) f(\sqrt{\tau}(\alpha x - \theta) + F^{-1}(x) + \sqrt{\tau}\underline{\theta}_{j-1}) d\theta}.$$

¹⁴Given the strategic complementarity, it follows from the Milgrom and Roberts (1990) argument that the unique cutoff equilibrium in Lemma 1 means it is the unique rationalizable strategy. However, under the endogenous information disclosure, if the agents in the past were attacking more aggressively, then surviving those attacks means the regime is likely to be strong. This, in turn, gives an incentive to the agents to not attack, and thus, the strategic complementarity may be violated. Therefore, the same generalization to rationalizable strategy may not hold for Lemma 3, and the existence of nonmonotone equilibria is an open question for future research.

¹⁵Note that $\theta \geq A_j$ is sufficient for the regime to pass the $(j+1)$ th viability test. It is not necessary. If $\theta = A_j > \underline{\theta}_{j-1}$, then the aggregate past attack is strictly lower than $\underline{\theta}_{j-1}$ (unless $\underline{\theta}_{j-1} = 0$). Thus, at $\theta = A_j$, the aggregate attack is less than the right-hand side of the equation (A_j^α) .

Note that if $\underline{\theta}_{j-1} = 0$ (as in Lemma 2), then this belief is equal to $G(x, \alpha)$ as defined in equation (G). Since π and f are bounded away from 0, for $x > 0$, the belief of the marginal agent $\hat{s}_j$ is at least

$$1 - \frac{\int_{\underline{\theta}_{j-1}}^{\underline{\theta}_{j-1} + \alpha x} \bar{\pi} \bar{f} d\theta}{\int_{\underline{\theta}_{j-1}}^{\underline{\theta}_{j-1} + \alpha x + \frac{1}{\sqrt{\tau}} F^{-1}(x) + \frac{1}{2\sqrt{\tau}}} \underline{\pi} \underline{f} d\theta} = 1 - \left(\frac{\bar{\pi} \bar{f}}{\underline{\pi} \underline{f}} \right) \left(\frac{\alpha x}{\alpha x + \frac{1}{\sqrt{\tau}} \left(F^{-1}(x) + \frac{1}{2} \right)} \right) \\ = \underline{G}(x, \alpha),$$

while for $x = 0$, the belief of the marginal agent $\hat{s}_j$ is 1.

Therefore, for $\alpha < \hat{\alpha}$ (as in Corollary 1), given M_j is true and given any possible fundamental cutoff $\underline{\theta}_{j-1}$, for any cutoff strategy $\hat{s}_j$, the marginal agent with cutoff signal $\hat{s}_j$ believes that the regime survives (or $\theta \geq A_j(\hat{s}_j, \alpha, \underline{\theta}_{j-1})$) with probability strictly higher than p . This implies M_{j-1} is true. ■

Step C: Note that M_j is trivially true since there is no attack after the last group. If the mass of agents in each group satisfies $\alpha < \hat{\alpha}$, by induction, we can show that in any cutoff equilibrium, no agent will ever attack a viable regime. When the principal runs the viability tests with sufficient frequency, the agents are assured that no agent will attack when the regime passes the next viability test (M_j is true). This makes the current positive viability news persuasive (N_j is true). Thus, under a sufficiently diffused policy ($J > \hat{J}$), the coordination risk unravels from the end. This proves our main result in Theorem 1.

Frequent Viability Test versus Stress Test.—Inostroza and Pavan (2018) and Goldstein and Huang (2016) propose a one-time “stress test” policy: the regime passes the stress test only if $\theta \geq k$ for some k . The authors demonstrate that if the stress test is sufficiently tough $k > \hat{k}$, then no agent will attack the regime that passes the stress test even in the cutoff equilibrium where agents attack most aggressively. In contrast, our paper considers frequent tests with the weakest possible strength or frequent viability tests. Although a one-time viability test may not be sufficiently “tough,” when the principal runs the tests with sufficient frequency, it can dissuade agents from attacking a viable regime.

Any viable regime survives when viability tests are repeated with sufficient frequency. In contrast, under a one-time stress test with $k > \hat{k}$, a viable regime whose strength is below k fails. Thus, from an ex ante perspective, the principal benefits from adopting a sufficiently frequent viability test policy compared with a sufficiently tough one-time test policy when the agents move sequentially. Furthermore, unlike the viability test, the stress test policy violates ex post incentive compatibility; if the principal misreports that the regime has passed the test when $\theta \in [0, k)$, then the agents will not attack and the regime will survive. This is not so for viability tests; if the regime fails a viability test, then the principal cannot benefit from misreporting it.¹⁶

¹⁶Thus, implementing a stress test policy requires commitment power as is standard in the Bayesian persuasion literature. Implementing viability tests do not require such commitment power. However, if repeating the tests are

IV. Application: Asynchronous Debt Structure

Consider a standard maturity mismatch problem, in which a borrower finances a long-term illiquid investment project by short-term debt contracts. Creditors may not roll over the short-term debts if they believe others will not do so. Hence, borrowers with reasonable sound fundamental may fail due to coordination failure among creditors in rolling over their debt. Any financial institution that performs liquidity transformation, e.g., commercial banks, hedge funds, and mutual funds, could be exposed to such panic-based runs. Brunnermeier (2009) argues that this coordination risk resulting from the maturity mismatch is a primary cause of the recent financial crisis.

Asynchronous debt structures that prevent creditors from withdrawing simultaneously are fairly common (see He and Xiong 2012). Choi, Hackbarth, and Zechner (2018) provides empirical evidence that corporate bond issuers diversify debt roll-overs across dates; they find that there is a significant increase in maturity dispersion for firms facing high rollover risk. Hedge funds and mutual funds are also allowed to lift redemption gates to limit momentary liquidity outflows.¹⁷ We are seeking a theoretical rationale behind the asynchronous debt structure.

Let us consider a stylized model in which there is a unit mass of creditors, and each of them has lent 1 unit to the borrower. The borrower uses this funding to finance some illiquid investment. An asynchronous debt structure J ensures only $\alpha = 1/J$ proportion of debt matures at time $(j - 1)\alpha$ for $j = 1, 2, \dots, J$. If the borrower fails to service the debt at any of these maturity dates, she must go through the bankruptcy process, which is public information to all creditors. This enforces the disclosure of the borrower's viability at each maturity date. Hence, the debt structure J is equivalent to the policy J we introduced in Section I.¹⁸

The asynchronous debt structure separates creditors into groups according to dispersed maturity dates. Group j creditors refers to the creditors who decide whether to withdraw ($a_i = 1$), or roll over the debt ($a_i = 0$) at time $(j - 1)\alpha$. The return from the illiquid investment R is realized at time 1. The variable θ indicates the liquidity position of the borrower. The borrower can use the liquid assets and the access to the extra funding source to sustain total withdrawals up to θ .¹⁹

The information environment is assumed to be the same as that described in Section I. The borrower chooses the debt maturity structure J prior to θ being realized. Thus, the debt structure itself does not provide information about θ .²⁰ All creditors know the maturity date of their debt as well as the overall debt structure. Let w_j

costly, then this may not be the case. It is important that the agents do not doubt that the principal will stop conducting the tests in the future.

¹⁷The redemption gates policy forces investors to make withdrawals asynchronously. Hedge fund managers can lift investor-level gates to limit investors' redemptions within a certain period. A common investor-level gate limits redemptions to 25 percent of an investor's money each quarter over four quarters (see Alistair Barr, "Hedge Funds Try New Way to Avoid Big Redemptions," *MarketWatch*, June 10, 2010). On October 14, 2014, SEC Rule 2a-7 was amended to enable managers of Money Market Mutual Funds to set redemption gates within a certain period during which a fund's liquidity position is unfavorable.

¹⁸We do not model the ex ante lending decision; rather, we focus on analyzing the rollover problem. However, ceteris paribus, if the borrower is less likely to default, then the creditors will be more willing to lend in the first place. Thus, such modification will only reinforce our result.

¹⁹The determination of the borrower's liquidity position θ is outside of the model. Note, θ may be negative if the liquidity outflow is high, due to the borrower's business operations and/or derivative positions, for example.

²⁰If the borrower only undertakes such a debt structure in times of distress, then this may defeat the purpose of the policy.

be the proportion of group j creditors who decide to withdraw. Given the debt structure J , group j creditors know whether the borrower is still viable at time $(j-1)\alpha$, i.e., $\theta \geq \alpha \sum_{l=1}^{j-1} w_l$.

If the borrower cannot service the withdrawal before time $(j-1)\alpha$, she defaults and all creditors whose debt matures after that time have no chance to make the withdrawal and obtain 0. Below, we describe the payoff for group j creditors when the borrower is still viable at time $(j-1)\alpha$.

By rolling over the debt, a group j creditor gets $1+r < R$ if the borrower can service all the withdrawals up to time 1; otherwise, the borrower obtains 0:

$$u(0, \theta, \{w_l\}_{l=1}^J) = \begin{cases} 1+r & \text{if } \theta \geq \alpha \sum_{l=1}^J w_l \\ 0 & \text{if } \theta < \alpha \sum_{l=1}^J w_l \end{cases}.$$

Upon withdrawal, the creditor gets his principal back at time $(j-1)\alpha$ if the borrower can sustain the withdrawal from group j , i.e., $\theta \geq \alpha \sum_{l=1}^j w_l$. Otherwise, the borrower defaults instantly, and the withdrawing creditors in group j split the remaining liquid assets θ .²¹ Thus, the payoff for group j creditors from withdrawal is

$$u(1, \theta, \{w_l\}_{l=1}^j) = \begin{cases} 1 & \text{if } \theta \geq \alpha \sum_{l=1}^j w_l \\ \frac{\theta - \alpha \sum_{l=1}^{j-1} w_l}{\alpha w_j} & \text{if } \theta < \alpha \sum_{l=1}^j w_l \end{cases}.$$

This payoff specification is different from that which is introduced in Section I. The payoff from rolling over depends on whether the borrower can sustain all withdrawals to time 1, while that from withdrawing only depends on whether the borrower can withstand the withdrawal from the current group. This asymmetry in payoff may increase the incentive for creditors to withdraw. Moreover, unlike in Section I, the payoff from withdrawing is not constant when the borrower defaults. This payoff specification is similar to that presented in Goldstein and Pauzner (2005); their specification indicates that the payoff from withdrawing is negatively dependent on aggregate withdrawal.²²

The following proposition shows that, despite these differences, when the borrower adopts a sufficiently asynchronous debt structure, the public information of the borrower's continued viability, i.e., success in rolling over debt maturing early, can avert panic-based runs altogether.

PROPOSITION 1: *A sufficiently asynchronous debt structure eliminates the risk of panic-based debt runs.*

²¹ For simplicity, we assume that the long-term illiquid investment has no liquidation value before it matures and the borrower cannot borrow against the return R .

²² We adopt this payoff specification for simplicity. We prove our result for a more general payoff structure. One can easily accommodate any reasonable payoff specification for withdrawing and rolling over when the borrower defaults.

Our result is robust to the case when payoff from attacking (or withdrawal) is independent of future attacks because we prove our main result using the statement M_j inductively, i.e., no one withdraws if the borrower remains viable after group j creditors had made their decisions. Thus, the asymmetry disappears since under M_j ; so long as the borrower can service the current withdrawal, no future withdrawal will occur ($w_l = 0$ for all $l > j$). Accommodating a general payoff structure broadens the scope of applying our proposed policy. Although the payoffs depend on both the fundamental and withdrawals, we can always find a threshold $\bar{p} < 1$ such that a creditor never withdraws if he believes that probability of no default is higher than $\bar{p}$ so long as the payoffs are bounded. Thus, a sufficiently asynchronous debt structure with $J > 1/\hat{\alpha}(\bar{p}, \tau)$ eliminates the chance of panic-based runs.

V. Discussion

Applying our theory to debt run provides a rationale for why borrowers often adopt asynchronous debt structures in practice. The insight also applies to regime change game in general (under more general payoff structure than that specified in Section I). In this section, we consider some variations of our setup to understand what could make it more or less difficult to dissuade agents from attacking. Moreover, to identify the essential features of the setup without which such persuasion may not be possible. First, let us consider the policy in which the principal discloses all information to the agents.

A. Full Disclosure

Suppose that the principal knows the underlying fundamental and discloses the true θ at the beginning. Furthermore, she discloses the information regarding past attacks with high frequency. Under such environment, when $\theta < 1$, it is easy to see that there exists an equilibrium in which all agents attack regardless of the history. Hence, any regime with $\theta < 1$ cannot survive.

Note that in our setting, an agent cannot make a difference in the aggregate attack by unilateral deviation. If this is not the case, full disclosure will lead to a very different result. Let us assume that attacking is the payoff-dominant action, i.e., $b_1 < b_0$. This means the agents' incentives are not aligned with the principal, and this is possible in practice. Consider, for example, a currency attack game as that presented in Morris and Shin (1998): attacking is costly but would be profitable if agents can coordinate ($b_0 > 0, c_1 < 0$), while not attacking has zero payoff ($b_1 = c_0 = 0$). In this case, attacking the regime of fixed exchange rate is the payoff-dominant action, but the policy designer may have incentive to defend it.

Consider the following simple example. Two agents are moving sequentially, one agent is equivalent to half mass, and the principal discloses all the information about θ and the past attack, regardless of whether the first agent has attacked. Suppose that $\theta \in (1/2, 1)$.²³ Then, the second agent will attack if the first agent attacks ($b_0 > b_1$) and will not attack if the first agent does not attack ($b_1 > c_1$). Since

²³The argument for $\theta < 1/2$ is straightforward because attacking is the dominant action for the first agent.

attacking is the payoff-dominant action ($b_0 > b_1$), the first agent attacks. Thus, it follows from backward induction that both agents will take the payoff-dominant action: attack. This simple example demonstrates that full disclosure may fail to dissuade agents from attacking. Thus, even if full disclosure is a feasible policy, the principal may not want to adopt such a policy.

The result that agents play the payoff-dominant action even generalizes to repeated games. In contrast with the folk theorem results, Dutta (2012) (see also Lagunoff and Matsui 1997) shows that if this coordination game is repeated (finitely but sufficiently many times), and agents alternately get chances to revise their decisions, then the agents will soon move to playing the payoff-dominant action regardless of the current state.

In contrast, our selective partial disclosure policy is persuasive regardless of whether attacking is the payoff-dominant action. If attacking is the payoff-dominant action, then it may be more difficult to dissuade the agents from attacking. Let us now move to the comparative statics and examine the factors that makes persuasion more difficult.

B. Comparative Statics

We have shown that a sufficiently diffused policy is persuasive. However, whether a diffused policy J qualifies as a sufficiently diffused policy depends on the model parameters. In our simple model, two parameters in which we are interested are p (reluctance) and τ (precision). We may expect that if the agents are more reluctant to follow when the principal recommends that they do not attack, it will be more difficult to persuade them. Furthermore, if the agents have more precise signals, it will be more difficult to persuade them to ignore their private signals.

It is easy to see that $\underline{G}(x, \alpha; \tau)$ is decreasing in α and τ . This in turn implies that $\hat{\alpha}(p, \tau)$ is decreasing in p and τ .²⁴ However, this does not exactly capture that the principal must diffuse more when p or τ is higher. It is because $\hat{J}$ is only a sufficient threshold, i.e., even if $J < \hat{J}(p, \tau)$, the diffused policy can be persuasive. However, under the uniform prior, we can provide a tight bound, i.e., a necessary and sufficient threshold for J . The following proposition describes the comparative statics results under the uniform prior.

PROPOSITION 2: *Under the uniform prior, there exists $J^*(p, \tau)$ such that a diffused policy is persuasive if and only if $J > J^*(p, \tau)$, where*

(i) *for $p' \geq p$, $J^*(p', \tau) \geq J^*(p, \tau)$, and*

(ii) *for $\tau' \geq \tau$, $J^*(p, \tau') \geq J^*(p, \tau)$.²⁵*

The result that $J^*(p, \tau)$ is not only sufficient but also necessary, follows from the fact that under the uniform prior, the belief of the marginal agent in group j regarding

²⁴Consider, for example, $p' \geq p$. If $\alpha < \hat{\alpha}(p', \tau)$, then $\underline{G}(x, \alpha; \tau) > p' > p$ for all $x \in [0, 1]$. Hence, $\hat{\alpha}(p', \tau) \leq \hat{\alpha}$

the survival of the regime (as defined in (G)) is decreasing in the group size α for any given $\underline{\theta}_{j-1}$.

This implies that when p is higher, a smaller group size, or equivalently, higher diffusion is required to make $G(x, \alpha; \tau) > p$ for all $x \in [0, 1]$. If $b_1 < b_0$, then an agent gets a higher payoff from attacking a regime that does not survive, as opposed to not attacking a regime that does survive. Thus, attacking is the payoff-dominant action, and the principal's interest is not aligned with the agents (recall the currency attack example). If attacking is the payoff-dominant action ($b_1 < b_0$), then p is higher (see equation (2)). From the proposition above, it follows that it will be harder to dissuade the agents from attacking. Moreover, under the uniform prior, $G(x, \alpha; \tau)$ is also increasing in τ . Therefore, persuading the agents with more precise private signals to ignore their private information and follow the principal's recommendation requires higher diffusion.

In our basic setup, we consider a regime change game in which agents have noisy private information as is standard in global game literature. Suppose that the noisy information was public; thus, agents share a homogeneous belief over θ . Below, we argue that in such an environment, it is possible that no diffused policy can dissuade the agents from attacking.

C. Homogeneous Belief

Consider the same setup as in Section I with the following modification. The agents receive a noisy public signal $s = \theta + \sigma\epsilon$ instead of private signals. For simplicity, let us assume the prior is uniform. The principal observes the public signal and announces the policy after learning s . Conditional on the realized signal s , the agents share a homogeneous belief over θ . Under a homogeneous belief, the only way to dissuade agents from attacking is to convince them “even if others attack, the regime is very likely to survive.” Otherwise, there is always a possible equilibrium

believe $\theta \geq 1$, or even if others do not believe so, they may entertain the possibility that others believe $\theta \geq 1$ and so on. Thus, the principal may not be able to convince an agent who receives a low private signal such as $s_i = -\sigma/2$ that “even if others attack the regime is very likely to survive.” However, the principal may still be able to convince the agent that “others are not likely to attack.” Thus, the regime is likely to survive the attacks.

PROPOSITION 3: *Under the information environment with the uniform prior and noisy public signal,*

- (i) *a diffused policy can dissuade the agents from attacking only if the noisy public signal s is sufficiently high such that $\Pr(\theta \geq 1|s) > 0$, and the required diffusion $J^*(s)$ increases when s decreases.*
- (ii) *If s is sufficiently small such that $\Pr(\theta \geq 1|s) = 0$, then no diffused policy can dissuade the agents from attacking.*
- (iii) *From an ex ante perspective, for any $\theta < 1$, there is a positive probability that (ii) will happen.²⁶*

The results (i) and (ii) follow from the argument preceding the proposition. For (iii), note that when $\theta < 1$, there is always a positive probability that the signal is so low that it becomes a common belief that $\theta < 1$; thus, no diffused policy can dissuade the agents from attacking. There is always a possible equilibrium in which all agents attack. In this sense, a diffused policy cannot eliminate the coordination risk when agents have noisy public information. This shows that private information environment is essential for our result.

VI. Conclusion

This paper proposes a simple policy called frequent viability tests that may be used to eliminate the strategic uncertainty in a global game of regime change. The frequent viability tests diffuse the coordination risk that would otherwise be concentrated at one point in time. We show that when the principal sufficiently diffuses the coordination risk, agents ignore their private information and follow the principal’s recommendation. The underlying mechanism is as follows. When the principal recommends the agent not to attack a regime that passes the viability test, and when the agents are assured that the agents in the subsequent groups will follow the principal’s recommendation (statement M_j), then a sufficiently small group of agents will follow the recommendation as well (argument N_j).

This result contributes to dynamic information design literature. From a methodological perspective, our paper develops an inductive argument to show that the

²⁶ Under unbounded noise (for example, the Gaussian distribution), the impossibility in (ii) and (iii) go away. However, for any J (however large), we can always find s such that $J^*(s) > J$. In this sense, although a diffused policy can work for any possible realization of public signal, the required diffusion can be unrealistically high.

coordination risk unravels from the end. From an applied perspective, we show that a sufficiently asynchronous debt structure eliminates the possibility of panic.

Readers may wonder that, since the agents move sequentially, some exogenous shock may hit the fundamental or new information may arrive over time. In the online Appendix, we detail the robustness of our main result to such perturbations in the basic model. Within the scope of this paper, we refrain from discussing the effectiveness of limited diffusion (when sufficient diffusion is not feasible) or optimal information disclosure under endogenous timing of attack. We believe that these are promising directions for future research.

APPENDIX

PROOF OF CLAIM 1:

First, we show that $\exists x \in [0, 1]$ such that $G(x, \alpha^*) \leq p$. Otherwise, it follows from uniform continuity that $\exists \delta > 0$ such that $\forall \alpha \in [\alpha^*, \alpha^* + \delta]$, $\forall x \in [0, 1]$, $G(x, \alpha) > p$. This implies $\alpha^* + \delta \in \mathcal{P}$, which contradicts the fact that $\alpha^* = \sup \mathcal{P}$.

Consider the nontrivial case in which $\alpha^* < 1$. For any $\alpha^{**} > \alpha^*$, by definition of α^* , $\alpha^{**} \notin \mathcal{P}$. Therefore, there exist $x^0 \in [0, 1]$ and $\alpha^0 \in (0, \alpha^{**}]$ such that $G(x^0, \alpha^0) \leq p$. Let $A^0 = x^0 \cdot \alpha^0 \in [0, \alpha^0]$. Consider any $\alpha^1 \geq \alpha^0$ and take $x^1 = A^0/\alpha^1$. Then, $x^1 \leq x^0$. Let us define the cutoff strategies $\hat{s}^1$ and $\hat{s}^0$ corresponding to x^1 and x^0 as in the aggregate condition (A^α) . Then,

$$\hat{s}^1 = \frac{1}{\sqrt{\tau}} F^{-1}(x^1) + A^0 \leq \frac{1}{\sqrt{\tau}} F^{-1}(x^0) + A^0 = \hat{s}^0.$$

Therefore, it follows from log-concavity of f that

$$G\left(\frac{A^0}{\alpha^1}, \alpha^1\right) = P(\theta \geq A^0 | \hat{s}^1, \theta \geq 0) \leq P(\theta \geq A^0 | \hat{s}^0, \theta \geq 0) = G\left(\frac{A^0}{\alpha^0}, \alpha^0\right).$$

Thus, for any $\alpha^1 \geq \alpha^0$, there exists $x^1 \in [0, 1]$ such that $G(x^1, \alpha^1) \leq G(x^0, \alpha^0) \leq p$. In particular, for $\alpha^1 = \alpha^{**}$, there is some $x \in [0, 1]$ such that $G(x, \alpha^{**}) \leq p$. ■

PROOF OF PROPOSITION 1:

Let us consider a general payoff structure for group j agents given the borrower has not failed yet, i.e., $\theta \geq \alpha \sum_{l=1}^{j-1} w_l$, as follows:

$$u(0, \theta, \{w_l\}_{l=1}^J) = \begin{cases} b_1(\theta, \{w_l\}_{l=1}^J) & \text{if } \theta \geq \alpha \sum_{l=1}^J w_l \\ c_0(\theta, \{w_l\}_{l=1}^J) & \text{if } \theta < \alpha \sum_{l=1}^J w_l \end{cases};$$

$$u(1, \theta, \{w_l\}_{l=1}^J) = \begin{cases} c_1(\theta, \{w_l\}_{l=1}^J) & \text{if } \theta \geq \alpha \sum_{l=1}^j w_l \\ b_0(\theta, \{w_l\}_{l=1}^J) & \text{if } \theta < \alpha \sum_{l=1}^j w_l \end{cases}.$$

Let us define the net payoff from rolling over as opposed to withdrawal as $\bar{u}(\theta, \{\mathbb{w}_l\}_{l=1}^J) := b_1(\theta, \{\mathbb{w}_l\}_{l=1}^J) - c_1(\theta, \{\mathbb{w}_l\}_{l=1}^J)$ when the borrower remains viable until the end, and as $\underline{u}(\theta, \{\mathbb{w}_l\}_{l=1}^J) := c_0(\theta, \{\mathbb{w}_l\}_{l=1}^J) - b_0(\theta, \{\mathbb{w}_l\}_{l=1}^J)$ when the borrower fails at time $j\alpha$ or after it.

ASSUMPTION 1: *The payoff has the following properties:*

- (i) (Complementarity) $\bar{u}(\cdot) \geq 0$ and $\underline{u}(\cdot) \leq 0$.
- (ii) (Boundedness) *There exist some finite numbers $\underline{m}$, $\bar{m}$, $\underline{n}$ and $\bar{n}$ such that $0 < \underline{n} \leq \bar{u}(\cdot) \leq \bar{n}$ and $0 < \underline{m} \leq -\underline{u}(\cdot) \leq \bar{m}$.*

The first part of the assumptions says that if the borrower is going to remain viable till time 1, then the agent is better off by rolling over, and if the borrower cannot withstand the withdrawal from group j , then the agent is better off by withdrawing. This captures the strategic complementarity. The second part of the assumption says that these net payoffs are bounded.

The payoff specification in the debt run application is a special case of this general payoff structure with $\bar{u} = r > 0$ and $-\underline{u} \in [0, 1]$. Below, we show that under the general payoff structure that satisfy Assumption 1, the argument N_j holds true for any $j = 1, 2, \dots, J$.

In this application, statement M_j can be interpreted as the following: no creditor from later groups will withdraw if the borrower can service the withdrawal from group j , i.e., $\mathbb{w}_l = 0$ for all $l > j$. That means $\theta \geq \alpha \sum_{l=1}^J \mathbb{w}_l$ is equivalent to $\theta \geq \alpha \sum_{l=1}^j \mathbb{w}_l$. Hence, we will follow the proof of Lemma 3 to prove the argument N_j .

Consider the marginal agent who receives the cutoff signal $\hat{s}_j$. Upon receiving the public news that the borrower is still viable at $(j - 1)\alpha$, he believes $\theta \geq \underline{\theta}_{j-1}$ for some $\underline{\theta}_{j-1}$. He also understands that the borrower will not default if $\theta \geq A_j(\hat{s}_j, \alpha, \underline{\theta}_{j-1})$ (as defined in (A_j^α)). Therefore, the marginal agent will roll over if he believes that the probability of no default

$$\Pr(\theta \geq A_j(\hat{s}_j, \alpha, \underline{\theta}_{j-1}) \mid \theta \geq \underline{\theta}_{j-1}) > \frac{1}{1 + E(\bar{u} \mid \theta \geq \underline{\theta}_{j-1})}$$

Hence, whenever $\alpha < \alpha^*(p, \tau)$, $G^u(x, \alpha; \tau) > p$ for all $x \in [0, 1]$ (regardless of $\underline{\theta}_{j-1}$). That means argument N_j is true for all $j \geq 1$ when $\alpha < \alpha^*(p, \tau)$.

On the other hand, for $\alpha \geq \alpha^*(p, \tau)$, there exists

in group $J - 1$ can also attack if they think the agents in group J will and so on. Thus, all attack is a possible equilibrium outcome.

- (iii) For any $\theta < 1$, there is a positive probability that a public signal $s \leq 1 - (\sigma/2)$ will be realized. Upon receiving such signal, it is possible that all the agents attack, regardless of the policy J . ■

REFERENCES

- 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): 711–56.
- Banerjee, Abhijit V.** 1992. "A Simple Model of Herd Behavior." *Quarterly Journal of Economics* 107 (3): 797–817.
- Bergemann, Dirk, and Stephen Morris.** 2013. "Robust Predictions in Games with Incomplete Information." *Econometrica* 81 (4): 1251–1308.
- Bergemann, Dirk, and Stephen Morris.** 2019. "Information Design: A Unified Perspective." *Journal of Economic Literature* 57 (1): 44–95.
- Bikhchandani, Sushil, David Hirshleifer, and Ivo Welch.** 1992. "A Theory of Fads, Fashion, Custom, and Cultural Change as Informational Cascades." *Journal of Political Economy* 100 (5): 992–1026.
- Brunnermeier, Markus K.** 2009. "Deciphering the Liquidity and Credit Crunch 2007–2008." *Journal of Economic Perspectives* 23 (1): 77–100.
- Carlsson, Hans, and Eric van Damme.** 1993. "Global Games and Equilibrium Selection." *Econometrica* 61 (5): 989–1018.
- Choi, Jaewon, Dirk Hackbarth, and Josef Zechner.** 2018. "Corporate Debt Maturity Profiles." *Journal of Financial Economics* 130 (3): 484–502.

- Milgrom, Paul, and John Roberts.** 1990. "Rationalizability, Learning, and Equilibrium in Games with Strategic Complementarities." *Econometrica* 58 (6): 1255–77.
- Morris, Stephen, and Hyun Song Shin.** 1998. "Unique Equilibrium in a Model of Self-Fulfilling Currency Attacks." *American Economic Review* 88 (3): 587–97.
- Morris, Stephen, and Hyun Song Shin.** 2003. "Global Games: Theory and Applications." In *Advances in Economics and Econometrics: Theory and Applications, Eighth World Congress*, Vol. 1, edited by Mathias Dewatripont, Lars Peters Hansen, and Stephen J. Turnovsky, 56–114. Cambridge: Cambridge University Press.
- Sákovics, József, and Jakub Steiner.** 2012. "Who Matters in Coordination Problems?" *American Economic Review* 102 (7): 3439–61.
- Szkup, Michal, and Isabel Trevino.** 2015. "Information Acquisition in Global Games of Regime Change." *Journal of Economic Theory* 160: 387–428.
- Varian, Hal R.** 1994. "Sequential Contributions to Public Goods." *Journal of Public Economics* 53 (2): 165–86.
- Vives, Xavier.** 2014. "Strategic Complementarity, Fragility, and Regulation." *Review of Financial Studies* 27 (12): 3547–92.