

Persuasion in Global Games

with Application to Stress Testing*

Nicolas Inostroza
Northwestern University

Alessandro Pavan
Northwestern University and CEPR

December 12, 2019

Abstract

We study robust/adversarial information design in global games of regime change. We show that the optimal policy coordinates all market participants on the same course of action. Importantly, while it removes any “strategic uncertainty,” it preserves heterogeneity in “structural uncertainty”. When the designer is constrained to public disclosures, we identify conditions under which the optimal policy is a “pass/fail” test, as well as conditions under which the test is monotone in the banks’ fundamentals. Finally, we show that the benefits from discriminatory disclosures come from “dividing-and-conquering” the market, and relate them to the type of securities issued by the banks.

JEL classification: D83, G28, G33.

Keywords: Global Games, Bayesian Persuasion, Information Design, Stress Tests.

*Emails: nicolasinostroza2018@u.northwestern.edu, alepavan@northwestern.edu. For comments and useful suggestions, we thank Marios Angeletos, Gadi Barlevy, Eddie Dekel, Tommaso Denti, Laura Doval, Stephen Morris, Marciano Siniscalchi, Bruno Strulovici, Jean Tirole, Xavier Vives, Nora Wagner, and seminar participants at various conferences and institutions where the paper was presented. Pavan acknowledges financial support from the NSF under the grant SES-1730483 and thanks Bocconi University for its hospitality during the 2017-2018 academic year. The usual disclaimer applies.

1 Introduction

Differing opinions on how stress tests should be undertaken are welcome and important...We need to move away from simple pass/fail policies (Piers Haben, Director, European Banking Authority, Financial Times, August 1, 2016).

Coordination plays a major role in many socio-economic environments. The damages to society of mis-coordination can be severe and often call for government intervention. Think of a major financial institution such as MPS (Monte dei Paschi di Siena, the oldest bank on the planet and the Italian third largest) trying to convince its major investors to keep pledging, despite growing rumors about the size of the bank’s non-performing loans. Financial analysts and market participants expect a default by MPS to trigger a sequence of “domino effects,” leading to a collapse in financial markets, and ultimately to a deep recession in the Eurozone and beyond (The Economist, July 7, 2016).

Confronted with such prospects, governments and supervising authorities have incentives to intervene. However, a government’s ability to calm the market by injecting liquidity into a troubled bank can be limited. For example, in Europe, legislation passed in 2015 prevents Eurozone member states from rescuing banks by purchasing assets or, more generally, by acting on the banks’ balance sheets. In such situations, the instrument of last resort often takes the form of interventions aimed at influencing market beliefs, for example through the design of stress tests, or other targeted information disclosures. The questions regulatory authorities face in designing such tests are the following: (a) What disclosures minimize the risk of market coordination failures? (b) Should all the information collected through the stress tests be passed on to the market, or should the supervising authorities commit to coarser policies, for example, simple announcements of whether or not the banks under scrutiny passed the tests? (c) Should the relevant authorities in charge be specific about the level or recapitalization asked to the banks, or simply announce that certain banks need further recapitalization, leaving it to the market to figure out the details? (d) Are there benefits to discriminatory disclosures, whereby different information is disclosed to different groups of investors?

In this paper, we develop a framework that permits us to investigate the above questions. We study the design of optimal information disclosures in markets in which a large number of agents (e.g., market investors) must choose whether to play a “socially desirable” action (e.g., roll over their loans), or “speculate” against a status quo regime (e.g., pull the money out of a potentially solvent but illiquid troubled bank). Market participants are endowed with heterogenous private information about relevant economic fundamentals, such as the size of a bank’s non-performing loans, or other elements of the bank’s balance sheet not in the public domain. A cash-constrained policy maker (e.g., a benevolent government, or a supervising authority such as the Federal Reserve Bank, or the European Banking Authority) can act so as to influence the market’s beliefs (for example, by designing a stress test), but is severely limited in its ability to use financial instruments to influence directly the market outcome.

While motivated by the design of stress tests, we abstain from many institutional details and, instead, cast the analysis in a broader class of games of regime change that can be used to shed light on similar questions also in other applications.¹ For example, in the context of currency crises, the policy maker may represent a central bank attempting to convince speculators to refrain from short-selling the domestic currency by releasing information about the bank’s reserves and/or about domestic economic fundamentals. Alternatively, the policy maker may represent the owners of an intellectual property, or more broadly the sponsors of an idea, choosing among different certifiers in the attempt to persuade heterogenous market users (buyers, developers, or other technology adopters) of the merits of a new product, as in Lerner and Tirole (2006)’s analysis of forum shopping.

The key novelty relative to the rest of the persuasion literature is that we explicitly account for the role that coordination plays among the receivers.² Furthermore, the latter are allowed to possess heterogenous private information prior to receiving additional information from the designer. At the theoretical level, these properties imply that, to derive the optimal persuasion strategy, one needs to study the effects of information disclosure not just on the agents’ first-order beliefs, but also on their *higher-order beliefs* (that is, the agents’ beliefs about other agents’ beliefs, their beliefs about other agents’ beliefs about their own beliefs, and so on). Equivalently, the optimal policy must be derived by accounting for how different information disclosures affect both the agents’ *structural uncertainty* (i.e., their beliefs about the underlying economic fundamentals), as well as the agents’ *strategic uncertainty* (i.e., the agents’ beliefs about other agents’ behavior).

The backbone of the analysis is a flexible global game of regime change in which, prior to receiving information from the policy maker (the information designer), each agent is endowed with an exogenous private signal about the strength of the regime (the critical size of attack above which the status quo collapses). In the absence of additional information, such a game admits a unique rationalizable strategy profile, whereby agents attack if, and only if, they assign sufficiently high probability to the underlying fundamentals being weak, and whereby regime change occurs only for sufficiently weak fundamentals.³

We take a “robust approach” to the design of the optimal information structure. We assume that, when multiple rationalizable strategy profiles are consistent with the disclosed information, the policy maker expects the agents to play according to the “most aggressive” strategy profile (the one that minimizes the policy maker’s payoff over the entire set of rationalizable profiles). This is an important departure from both the mechanism design and the persuasion literature, where the designer is typically assumed to be able to coordinate the market on her most preferred continuation equilibrium. Given the type of applications the analysis is meant for, such “robust approach” appears

¹For an account of key institutional details of the stress tests conducted in Europe, see, for example, Henry and Christoffer (2013) and Homar et al. (2016).

²See Bergemann and Morris (2017) for an excellent overview of this literature.

³Games of regime change have been used to model, among other things, currency crises, debt crises, political change, and standards adoption. See Morris and Shin (2006) for an excellent overview.

more appropriate.⁴

Our first result shows that the optimal policy has the “*perfect coordination property*.” It induces all market participants to take the same action, irrespective of the heterogeneity in the agents’ first- and higher-order (posterior) beliefs. In other words, the optimal policy completely removes any strategic uncertainty, while retaining heterogeneity in structural uncertainty. Under the optimal policy, each agent is able to predict the actions of any other agent, but not the beliefs that rationalize such actions. In particular, an agent who expects all other agents to refrain from attacking need not be able to predict whether most other agents do so because it is dominant for them not to attack or simply because they expect others to refrain from attacking. Such residual heterogeneity in structural uncertainty is key to minimizing the probability of regime change, irrespective of the selection of the strategy profile.

The above result is true irrespective of whether the policy maker is constrained to disclose the same information to all market participants or can engage in discriminatory disclosures, whereby different information is disclosed to different agents. The result is proved by showing that, starting from any initial collection of heterogeneous beliefs (formally, from any subset of the universal type space), and any (possibly discriminatory) disclosure policy, the policy maker can construct another policy that weakly improves upon the first one. The new policy is obtained from the original policy by augmenting the latter with a public announcement of whether or not the regime would have survived under the most aggressive rationalizable strategy profile consistent with the original policy. The new policy improves upon the original one because it induces all agents to refrain from attacking under all circumstances in which the regime would have survived under the original policy. The difficulty in establishing the result comes from the fact that the announcement that the regime would have survived under the most aggressive strategy profile consistent with the original policy carries information not just about the underlying fundamentals but also about other agents’ first- and higher-order beliefs (formally, it moves the agents’ beliefs into a new subset of the universal type space). We establish the result by showing that, under the new policy, at any step of the rationalizability process, agents are weakly less aggressive than under the original policy. That is, any agent who would have refrained from attacking under the original policy continues to do so under the new policy. In turn, this implies, that, in the limit, under the unique rationalizable profile, no agent attacks when hearing that the regime would have survived under the original policy.⁵

⁴If the designer could choose the rationalizable profile, she would fully disclose the fundamentals, and then recommend that all agents refrain from attacking, unless the regime is bound to collapse irrespective of the agents’ behavior. This is both uninteresting and unrealistic.

⁵Guaranteeing that no agent attacks under the strategy profile most advantageous to the designer is trivial. The difficulty is in showing that the same property holds across all possible rationalizable profiles. In the Supplementary Material, we show that the optimality of disclosure policies satisfying the perfect coordination property is a general feature of a large class of supermodular games with binary aggregate outcomes (e.g., games of regime change). In particular, the result extends to settings with an arbitrary number of agents with heterogeneous payoffs and beliefs that need not be consistent with a common prior, as well as to settings in which the policy maker can condition the

The implication of the above result for stress test design is that, in general, such tests should not be expected to generate consensus among market participants about the soundness of the financial institutions under scrutiny. Preserving, and, when possible, even enhancing the dispersion of beliefs is instrumental to a successful recapitalization of the banks. At the same time, there is no value in creating ambiguity about the response of the market.

In many cases of interest, we expect the policy maker to be unable to disclose different information to different market participants. Our second result pertains to such cases. It identifies primitive conditions under which the optimal non-discriminatory policy takes the form of a simple “pass/fail” test, with no further information disclosed to the market. We show that the optimality of such simple policies hinges on the policy maker expecting the agents’ beliefs to satisfy the following property: each agent believes that states of Nature in which the regime’s fundamentals are strong (i.e., in which the amount of non-performing loans is relatively small) are also states in which most agents expect the fundamentals to be strong.⁶ This property may be reasonable in many cases of interest and is consistent with what is typically assumed in the literature on coordination under incomplete information. Importantly, when such a property is not satisfied, the policy maker may be strictly better off disclosing information to the agents in addition to the announcement of whether or not the bank passed the test.

Our third result is about the optimality of deterministic monotone tests that pass with certainty all institutions whose fundamentals are strong and fail, with certainty, all institutions whose fundamentals are weak. We show that such policies are typically suboptimal. We identify sharp conditions under which such monotone rules are optimal and show that they are quite stringent. In persuasion environments with a single receiver, the optimality of such monotone policies is guaranteed by the supermodularity of the sender’s and the receiver’s payoffs (see Mensch (2017)). This is not the case with multiple receivers. For such policies to be optimal, the benefit the policy maker derives from avoiding default must grow with the strength of the underlying fundamentals sufficiently faster than the loss that each agent who refrains from attacking experiences in case of default.⁷ It seems hard to argue that such condition should be expected to hold in most cases of interest. Importantly, we show that the condition is sharp, in the sense that, when violated, one can construct non-monotone policies that improve upon the monotone ones. Of course, the political viability of such non-monotone policies is questionable, as most authorities may feel uncomfortable failing institutions with strong fundamentals while passing others with weaker fundamentals.

information disclosed to each agent on the agent’s prior beliefs. It also extends to economies where agents are boundedly rational in the sense of being able to perform only finitely many iterations in the computation of best responses (level-k reasoning).

⁶Formally, when the agents’ beliefs are parametrized by a uni-dimensional signal, this amounts to assuming that the signal distribution is log-supermodular or, equivalently, satisfies the monotone likelihood ratio property.

⁷See also Goldstein and Leitner (2017) for a similar point. In their environment, the optimality of non-monotone policies stems from the possibility of implementing superior risk sharing agreements. In our environment, instead, from the fact that the bank’s investors play a coordination game under asymmetric information.

The above results contribute to the recent debate about the (sub)optimality of the stress tests conducted in the Eurozone after the sovereign-debt crises. Such tests have been criticized for being too vague about the precise levels of recapitalization asked to the banks and, more generally, for not disclosing the details of the various simulations run to determine the banks' performance in the most adversarial scenarios (see, e.g., the article "Stress tests do little to restore faith in European banks," Financial Times, August 1, 2017). Our results indicate that simple pass/fail policies (whereby authorities announce whether the financial institutions under scrutiny are expected to meet their financial obligations across a variety of adversarial scenarios, provided they raise capital according to a pre-specified recapitalization plan, but without getting into the details of the banks' balance sheet), need not be suboptimal. Importantly, optimal stress tests should be *transparent*, but not in the sense of revealing all the information gathered during the tests, but in the sense of facilitating coordination among the relevant investors.

The last two results about the optimality of simple pass/fail policies and of monotone tests pertain to situations in which the policy maker is constrained to disclose the same information to all market participants, which, empirically, appears the most relevant case. In Section 5, we come back to discriminatory disclosures and illustrate why, when feasible, such disclosures may improve upon their non-discriminatory counterparts. We show that the benefits of discriminatory disclosures do not come from the possibility of targeting agents with different beliefs with different disclosures. They come primarily from the possibility of *dividing-and-conquering* the market by enhancing the dispersion of first- and higher-order beliefs so as to make it more difficult for each agent to predict what rationalizes other agents' behavior. In particular, discriminatory disclosures can strictly improve upon non-discriminatory ones even in environments in which market participants share the same beliefs prior to receiving information through the stress tests. The intuition is similar to the one in the contracting-with-externalities literature (see, e.g., Segal (2003)).⁸ With discriminatory disclosures, the policy maker makes it dominant for certain agents not to attack, and then leverages on such a property by making it *iteratively dominant* for all other agents to refrain from attacking.

We conclude with a few results illustrating how the (strict) optimality of discriminatory disclosures relates to the sensitivity of the agents' payoffs to the underlying fundamentals. To gain some tractability, we specialize the analysis to an environment with Gaussian beliefs, where both the agents' exogenous signals and the endogenous signals disclosed by the policy maker are normally distributed. We show that discriminatory policies strictly improve upon non-discriminatory ones only when the sensitivity of the agents' payoffs to the underlying fundamentals is stronger in case of regime change than when the regime survives. In the context of our leading application, the result implies that non-discriminatory policies are optimal when the securities issued by the banks resemble equity claims. When, instead, they resemble bonds, discriminatory disclosures are typically superior to non-discriminatory ones. Importantly, as anticipated above, irrespective of whether or

⁸See also Moriya and Yamashita (2017) for an analysis of the benefits of discriminatory disclosures in team-production problems.

not discriminatory disclosures dominate non-discriminatory ones, there is nothing to be gained by mis-coordinating the response of the market. Optimal discriminatory disclosures can thus be implemented by having the supervising authority publicly announce whether or not a bank passed the test and, in case it did, disclosing different summary statistics of the test results to different groups of investors.

The rest of the paper is organized as follows. Below, we wrap up the introduction with a brief review of the most pertinent literature. Section 2 presents the model. Section 3 establishes the optimality of the perfect coordination property. Section 4 studies properties of optimal non-discriminatory policies: It identifies conditions under which the optimal policy has a simple pass/fail structure, as well as conditions under which the optimal policy is monotone in the fundamentals. Section 5 illustrates the benefits of discriminatory disclosures. Section 6 concludes. Proofs omitted in the text are either in the Appendix at the end of the document or in the Supplementary Material.

(Most) pertinent literature. The paper is related to different strands of the literature. The first strand is the literature on *information design*. This literature traces back to Myerson (1986), who introduced the idea that, in a general class of multi-stage games of incomplete information, the designer can restrict attention to private incentive-compatible action recommendations to the agents. Important recent developments include Kamenica and Gentzkow (2011), [?](#), and Ely (2017). These papers consider persuasion with a single receiver. The case of multiple receivers is less studied. Calzolari and Pavan (2006a) consider an auction setting in which the sender is the initial owner of a good and where the different receivers are bidders in an upstream market who then resell in a downstream market (see also Dworzak (2016) for an analysis of persuasion in other mechanism design environments with aftermarkets).⁹ Alonso and Camara (2016a) and Bardhi and Guo (2017) consider persuasion in a voting context, whereas Mathevet et al. (2016) and Taneva (2016) study persuasion in more general multi-receiver settings. Importantly, these papers assume that the receivers are homogeneously informed (share a common prior) about the underlying payoff-relevant parameters. Persuasion with ex-ante heterogeneously informed receivers is examined in Bergemann and Morris (2016a), Bergemann and Morris (2016b), Kolotilin et al. (2017), Alonso and Camara (2016b), Chan et al. (2016), Che and Hörner (2017), Basak and Zhou (2017), and Doval and Ely (2017).¹⁰ Bergemann and Morris (2016a) and Bergemann and Morris (2016b) characterize the set of outcome distributions that can be sustained as Bayes-Nash equilibria under arbitrary information structures consistent with a given common prior. Alonso and Camara (2016b) study public persuasion in a setting with multiple receivers with heterogeneous priors. Kolotilin et al. (2017) consider a screening environment whereby the designer elicits the agents' private information prior to disclosing further information to them. Chan et al. (2016) study pivotal persuasion in a voting environment similar to

⁹Related is also Calzolari and Pavan (2006b). That paper studies information design in a model of sequential contracting with multiple principals.

¹⁰See also Shimoji (2017) and Arieli and Babichenko (2017). These papers, however, abstract from strategic interactions among the receivers.

the one in Alonso and Camara (2016a), but where the sender is allowed to communicate privately with the voters.¹¹ Che and Hörner (2017) consider dynamic disclosures in a social experimentation setting where the designer’s information depends on the agents’ past experimentation decisions. Basak and Zhou (2017) and Doval and Ely (2017) study dynamic games in which the designer can control both the agents’ information and the timing of their actions.

The present paper contributes to this literature by focusing on large markets where the receivers play a (global) coordination game under dispersed information. The approach is also different in that we assume the designer evaluates any disclosure rule on the basis of the least advantageous rationalizable strategy profile it induces.

The second strand is the literature on *stress test design* and *regulatory disclosures* in the financial system. For an excellent overview of this literature see, e.g., Goldstein and Sapra (2014).¹² Close in spirit is the work by Goldstein and Leitner (2015). That paper studies the design of stress tests by a regulator facing a competitive market, where agents hold homogeneous beliefs about the bank’s balance sheet.¹³ In contrast, in the present paper, we consider the design of stress tests by a policy maker facing a continuum of investors with heterogeneous private beliefs. We also model explicitly the coordination game among market participants. Bouvard et al. (2015) study a credit rollover setting similar to ours where a policy maker must choose between transparency (full disclosure) and opacity (no disclosure) but cannot commit to a disclosure policy. In contrast, we assume the policy maker can fully commit to her disclosure policy and allow for flexible information structures. Alvarez and Barlevy (2015) study the incentives of banks to disclose balance sheet (hard) information in a setting where the market is not able to observe how banks are exposed to each others’ risks.¹⁴ Orlov et al. (2017) consider the joint design of stress tests and capital requirements in a setting where multiple banks have correlated exposures to exogenous shocks.¹⁵ Related is also Goldstein and Huang (2016). That paper studies persuasion in a coordination setting similar to ours, but restricts the designer to announcing whether or not the fundamentals fall below a given threshold. We allow for flexible information structures, but also identify conditions under which monotone disclosures similar to those in Goldstein and Huang (2016) are optimal.

The present paper contributes to this literature in a few important ways: (a) it shows that optimal stress tests should not be expected to create conformism in investors’ beliefs about banks’

¹¹Discriminatory persuasion in a voting setting is examined in Wang (2015). That paper, however, restricts the sender to conditionally i.i.d. signals.

¹²See also Morgan et al. (2014), Flannery et al. (2017), and Petrella and Resti (2013). The first two papers provide evidence that the tests conducted in the US revealed information not already in the market system, whereas the second paper provides similar evidence for stress tests conducted in the EU.

¹³See also Williams (2017) for a related analysis of stress test design in a bank-run model à la Allen and Gale (1998), with homogenous investors.

¹⁴See also Corona et al. (2017) for an analysis of how stress tests disclosures may favor banks’ coordinated risk taking in the spirit of Farhi and Tirole (2012).

¹⁵See also Faria-e Castro et al. (2016) and Garcia and Panetti (2017) for a joint analysis of stress tests and government bailouts.

fundamentals but should be sufficiently transparent to eliminate any ambiguity about the market response to the tests; (b) it identifies conditions under which optimal stress tests take the form of simple pass/fail announcements; (c) it provides conditions for optimal tests to be monotone; and (d) it discusses benefits of discriminatory disclosures and relate them to the type of securities issued by the banks.

Finally, the paper is related to the literature on *global games with endogenous information*. Angeletos et al. (2006), and Angeletos and Pavan (2013) consider settings whereby a policy maker, endowed with private information, engages in costly actions to influence the agents' behavior. Edmond (2013) considers a similar setting but assumes the cost of policy interventions is zero and agents receive noisy signals of the policy maker's action. Angeletos et al. (2007) consider a dynamic model in which agents learn from the accumulation of private signals over time and from the (possibly noisy) observation of past outcomes. Cong et al. (2016) consider a dynamic setting similar to the one in Angeletos et al. (2007) but allowing for policy interventions. Denti (2015), Szkup and Trevino (2015), Yang (2015) and Morris and Yang (2016) consider global games where, prior to committing their actions, agents acquire information about payoff-relevant variables.

2 Model

Players and Actions. The economy is populated by a big player, the policy maker, who seeks to influence the fate of a regime, and a (measure-one) continuum of atomistic agents, who must choose whether or not to attack the regime. We index the agents by i and assume they are distributed uniformly over $[0, 1]$. We denote by $a_i = 1$ the decision by agent $i \in [0, 1]$ to attack, and by $a_i = 0$ the decision by the same agent to not attack. We then denote by $A \in [0, 1]$ the aggregate size of the attack.

Fundamentals. The payoff structure is parameterized by the random variable $\theta \in \mathbb{R}$. This variable parametrizes both the strength of the status quo (i.e., the critical size of the aggregate attack above which the status quo collapses) and the agents' preferences. We will refer to θ as the "underlying fundamentals." It is common knowledge that θ is drawn from an absolutely continuous distribution F , with a smooth density f strictly positive over $\mathbb{R}$, and first and second moments given by μ_θ and σ_θ^2 , respectively.

Exogenous information. Each agent $i \in [0, 1]$ is endowed with a noisy private signal $x_i \in \mathbb{R}$ about the underlying fundamentals. Conditional on θ , the signals x_i are i.i.d. draws from the cdf $P(x|\theta)$ with associated density $p(x|\theta)$. The cross-sectional distribution of exogenous signals in the population is denoted by $\mathbf{x} \in \mathbb{R}^{[0,1]}$, and, for any $\theta \in \Theta$, $\mathbf{X}(\theta) \subseteq \mathbb{R}^{[0,1]}$ denotes the collection of cross-sectional distributions of signals consistent with the fundamentals being equal to θ .

Regime outcome. Let $r \in \{0, 1\}$ denote the regime outcome. We denote by $r = 1$ the event that the status quo is abandoned, and by $r = 0$ the complement event in which the status quo is preserved. Regime change occurs, i.e., $r = 1$, if, and only if, $R(\theta, A) \leq 0$, where R is a continuous

function, strictly increasing in θ , and decreasing in A .

Dominance Regions. There exist thresholds $\underline{\theta}, \bar{\theta} \in \mathbb{R}$ such that $R(\underline{\theta}, 0) = R(\bar{\theta}, 1) = 0$. Irrespective of the size of the attack, the status quo collapses when $\theta < \underline{\theta}$, and survives when $\theta \geq \bar{\theta}$.

Payoffs. The policy maker's payoff is equal to

$$U^P(\theta, A) = \begin{cases} W(\theta, A) & \text{if } r = 0 \\ L(\theta) & \text{if } r = 1, \end{cases} \quad (1)$$

with the function W continuously differentiable and satisfying the following properties, for any $(\theta, A) \in \mathbb{R} \times [0, 1]$: (a) $W(\theta, A)$ is non-increasing in A ; (b) $W(\theta, A) - L(\theta) \geq 0$ if $R(\theta, A) > 0$. The first property says that, conditional on the status quo surviving the attack, the payoff to the policy maker decreases (weakly) with the size of the aggregate attack. The second property says that the policy maker would never prefer to see the status quo collapse when it survives.¹⁶ The assumption that, in case of regime change, L is invariant in A captures the idea that, when regime change occurs, the policy maker is indifferent as to the precise size of the attack that triggered regime change. For example, in the context of stress test design, conditional on default, the government is indifferent as to the level of speculation that led the financial institution into bankruptcy. All these assumptions are fairly standard in the literature on coordination under incomplete information (see, e.g., Angeletos and Pavan (2013) and the discussion therein).

The agents' payoff from attacking is normalized to zero, whereas their payoff from not attacking is equal to¹⁷

$$u(\theta, A) = \begin{cases} g(\theta, A) & \text{if } r = 0 \\ b(\theta, A) & \text{if } r = 1. \end{cases}$$

The functions g and b are continuously differentiable and satisfy the following assumptions, for any $(\theta, A) \in \mathbb{R} \times [0, 1]$:¹⁸ (a) $g_\theta(\theta, A), b_\theta(\theta, A) \geq 0$ and $g_A(\theta, A), b_A(\theta, A) \leq 0$; (b) $g(\theta, A) > 0 > b(\theta, A)$. In the context of stress-test design, the first assumption means that the payoff that an investor expects from pledging to a bank (weakly) increases with the bank's amount of performing loans (the fundamentals) and with the number of other investors who also pledge $(1 - A)$. The second assumption says that pledging yields a payoff higher than investing in other instruments in case default does not occur, whereas the opposite is true in case of default. These assumptions readily extend to other applications.

A simple structure often encountered in applications which is consistent with the assumptions above is one where W, L, g, b are all constants and where R is linear, i.e., $R(\theta, A) = \theta - A$. While

¹⁶This second property trivially holds when the fate of the regime is controlled directly by the policy maker, as in certain applications.

¹⁷In case of currency attacks and political change, it is customary to normalize the payoff from not attacking to 0. That is, to assume the "safe action" is not-attacking. This can be easily accommodated by changing the interpretation of the actions.

¹⁸The functions $g_\theta(\theta, A)$ and $b_\theta(\theta, A)$ are partial derivatives with respect to the θ dimension. Similarly, $g_A(\theta, A)$ and $b_A(\theta, A)$ are partial derivatives with respect to A .

all the results below extend to the more general payoff structure introduced above, the reader may focus on this simple structure to fix ideas.¹⁹

Disclosure Policies. The only instrument the policy maker possesses to influence the regime outcome is the design of a disclosure policy. Let $\mathcal{S}$ be a compact metric space defining the set of possible disclosures to the agents. Let $m : [0, 1] \rightarrow \mathcal{S}$ denote a *message function*, specifying, for each individual $i \in [0, 1]$, the endogenous signal $m_i \in \mathcal{S}$ disclosed to the individual. Let $M(\mathcal{S}) = \{\mathcal{S}^{[0,1]}\}$ denote the set of all possible message functions with codomain $\mathcal{S}$. A *disclosure policy* $\Gamma = (\mathcal{S}, \pi)$ consists of a set $\mathcal{S}$ along with a mapping $\pi : \Theta \rightarrow \Delta(M(\mathcal{S}))$ specifying, for each θ , a lottery whose realization yields the message function used to communicate with the agents. As is standard in the literature, the disclosure policy Γ itself does not convey any information about θ to the agents (in the context of stress test design, the assumption reflects the idea that the policy maker does not possess private information about the financial institutions under scrutiny prior to conducting the tests). Furthermore, the policy maker can credibly commit not to modify Γ once the latter is announced.²⁰

We restrict attention to policies Γ with the property that, when agents play according to the most aggressive rationalizable strategy profile consistent with Γ (formally defined below), the regime outcome is measurable in the policy maker’s information. This condition is satisfied, for example, by all policies that combine public signals of θ (of any nature) with private signals drawn independently across agents, given θ — see the Supplementary Material for a generalization.

Timing. The sequence of events is as follows:

1. The policy maker chooses a disclosure policy $\Gamma = (\mathcal{S}, \pi)$ and publicly announces it.
2. The fundamentals of the economy θ and the agents’ exogenous signals $\mathbf{x} \in \mathbf{X}(\theta)$ are realized.
3. A message function $m \in \text{supp}[\pi(\theta)]$ is drawn from the distribution $\pi(\theta) \in \Delta(M(\mathcal{S}))$ and used to disclose information to the agents.
4. Agents simultaneously choose whether or not to attack.
5. The regime outcome is determined by (θ, A) and payoffs are realized.

Adversarial/robust design. The policy maker evaluates any given policy Γ on the basis of the “worst outcome” consistent with the agents playing (interim correlated) rationalizable strategies. That is, given any policy Γ , the policy maker expects the market to play according to the “most aggressive rationalizable profile” defined as follows.

¹⁹The extra generality, however, plays a role for the results about monotone tests, as well as for the results relating the optimality of non-discriminatory policies to the type of securities issued by the banks.

²⁰See Leitner and Williams (2017) for a model in which the policy maker’s stress test protocol is her private information.

Definition 1. Given any policy Γ , the most aggressive rationalizable profile (MARP) consistent with Γ is the strategy profile $a^\Gamma \equiv (a_i^\Gamma)_{i \in [0,1]}$ that minimizes the policy maker's ex-ante expected payoff over all profiles surviving *iterated deletion of interim strictly dominated strategies* (henceforth IDISDS).

As it will become clear from the analysis below, the strategy profile a^Γ is, in fact, a Bayes-Nash equilibrium (BNE) of the continuation game that follows the announcement of Γ , and minimizes the policy maker's payoff state-by-state, and not just in expectation.

3 Perfect Coordination Property

We now turn to the first property of optimal policies.

Definition 2. A policy $\Gamma = (\mathcal{S}, \pi)$ satisfies the **perfect-coordination property** if, for any $\theta \in \Theta$, any distribution of exogenous information $\mathbf{x} \in \mathbf{X}(\theta)$, any message function $m \in \text{supp}[\pi(\theta)]$, any pair of individuals $i, j \in [0, 1]$, $a_i^\Gamma(x_i, m_i) = a_j^\Gamma(x_j, m_j)$, where $a^\Gamma = (a_i^\Gamma)_{i \in [0,1]}$ is the most aggressive rationalizable profile (MARP) consistent with the policy Γ .

Hence, a disclosure policy has the perfect-coordination property if it induces all market participants to take the same action, after any information it discloses. Now, for any $\theta \in \Theta$, any $\mathbf{x} \in \mathbf{X}(\theta)$, any $m \in \text{supp}[\pi(\theta)]$, let $r(\theta, \mathbf{x}, m; a^\Gamma) \in \{0, 1\}$ denote the regime outcome that prevails at θ when the agents receive exogenous information $\mathbf{x}$ and endogenous information m , and play according to MARP consistent with Γ .

Definition 3. The disclosure policy $\Gamma = (\mathcal{S}, \pi)$ is **regular** if, for any $\theta \in \Theta$, any $m \in \text{supp}[\pi(\theta)]$, any pair of exogenous signal distributions $\mathbf{x}, \mathbf{x}' \in \mathbf{X}(\theta)$, $r(\theta, \mathbf{x}, m; a^\Gamma) = r(\theta, \mathbf{x}', m; a^\Gamma)$.

Observe that, because exogenous signals x_i are i.i.d. draws from $P(x|\theta)$, when the policy Γ discloses the same information to all agents, the regime outcome under MARP is the same across any pair of signal distributions $\mathbf{x}, \mathbf{x}' \in \mathbf{X}(\theta)$ consistent with the fundamentals being equal to θ . The definition extends the same property to policies that disclose different information to different agents by requiring that, when the agents play according to MARP, the regime outcome remains measurable in the policy maker's information, that is, in (θ, m) . Also note that, implicit in the definition, is the requirement that MARP is well defined under Γ , which in turn requires the procedure of IDISDS in the continuation game that starts after the policy Γ is announced to be well defined. Hereafter, we denote by $r(\theta, m; a^\Gamma) \in \{0, 1\}$ the regime outcome that prevails at (θ, m) , when agents play according to MARP consistent with the (regular) policy Γ .

Theorem 1. *Given any regular policy Γ , there exists another regular policy Γ^* satisfying the **perfect coordination property** that yields the policy maker an expected payoff weakly higher than Γ .*

The formal proof is in the Appendix. Here we provide a heuristic sketch of the key arguments. The policy Γ^* is obtained from the original policy Γ by disclosing to the agents, in addition to the signals

disclosed by the original policy Γ , the regime outcome $r(\theta, m; a^\Gamma) \in \{0, 1\}$ that would have prevailed at (θ, m) under MARP consistent with the original policy Γ . The difficulty in establishing the result is that the agents' posterior beliefs about θ , with and without the extra information contained in $r(\theta, m; a^\Gamma)$, cannot be ranked (e.g., according to FOSD). Furthermore, the announcement that (θ, m) is such that the regime would have survived under a^Γ carries information not only about θ , but also about the distribution of first- and higher-order beliefs in the population. In principle, such extra information may permit the agents to coordinate on a strategy profile that is more aggressive than MARP under the original policy Γ .²¹

The key property that guarantees the optimality of policies satisfying the perfect coordination property is the “truncation” of beliefs induced by the policy Γ^* . The announcement that (θ, m) is such that regime change would not have occurred under MARP consistent with the original policy Γ makes the event $\{(\theta, m) \in \Theta \times M(\mathcal{S}) : m \in \text{supp}[\pi(\theta)], r(\theta, m; a^\Gamma) = 0\}$ common certainty among the agents. In addition, the new policy preserves the likelihood ratio of any two states

$$(\theta', m'), (\theta'', m'') \in \{(\theta, m) \in \Theta \times M(\mathcal{S}) : m \in \text{supp}[\pi(\theta)], r(\theta, m; a^\Gamma) = 0\}$$

for which regime change would not have occurred under MARP consistent with the original policy Γ . Leveraging on these properties, the proof in the Appendix shows that at any step in the procedure of iterated deletion of dominated strategies, any agent who would have refrained from attacking under the original policy Γ also refrains from attacking under the new policy Γ^* . The combination of the fact that the new policy Γ^* makes it common certainty among the agents that regime change would not have occurred under MARP consistent with the original policy Γ along with the fact that agents are weakly less aggressive under the new policy Γ^* than under the original policy Γ then guarantees that, under the new policy Γ^* , there is a unique rationalizable profile following the announcement that (θ, m) is such that $r(\theta, m; a^\Gamma) = 0$, and is such that all agents refrain from attacking.

Similarly, the public announcement that (θ, m) is such that regime change would have occurred under the original policy Γ (namely, that $r(\theta, m; a^\Gamma) = 1$) makes it common certainty among the agents that $\theta \leq \bar{\theta}$. Therefore the most aggressive rationalizable profile following such announcement features all agents attacking. That the new policy Γ^* improves upon the original policy Γ then follows from the fact that Γ^* maintains invariant the probability regime change occurs at any θ , while minimizing the size of the attack at each θ at which regime change does not occur.

The policy Γ^* thus removes any strategic uncertainty. When the signal $r = 0$ (alternatively, $r = 1$) is disclosed, each agent knows that all other agents will refrain from attacking (alternatively,

²¹That, under the new policy Γ^* , there exists a rationalizable profile that is outcome-equivalent to MARP under the original policy Γ is trivial and follows from argument similar to those establishing the Revelation Principle (see, e.g., Myerson (1986)). Under adversarial design, however, one needs to establish that the policy maker is better off under the new policy Γ^* , no matter the rationalizable profile that follows the announcement of Γ^* . The difficulty in establishing the result is thus akin to the difficulty in characterizing properties of optimal mechanisms in the *full implementation* literature where it is known that simple direct revelation mechanisms need not be optimal.

that they will attack), irrespective of their exogenous and endogenous private information (x_i, m_i) , and finds it optimal to do the same.

Importantly, while the policy Γ^* removes any strategic uncertainty, it preserves, and in some cases it even enhances heterogeneity in structural uncertainty. Under Γ^* , different agents holds different beliefs about the underlying fundamentals. Preserving heterogeneity in posterior beliefs about θ is key to the minimization of a risk of regime change. In fact, if agents knew the exact fundamentals, then, under the most aggressive rationalizable profile, they would all attack for any $\theta \leq \bar{\theta}$. More generally, if agents knew each others' beliefs, they could coordinate on a more aggressive continuation strategy profile inducing regime change over a larger set of fundamentals.

The policy Γ^* leverages on the fact that, when the public signal $r = 0$ is announced, agents remain uncertain as to whether other agents will refrain from attacking because they find it dominant to do so, or because they expect others to refrain from attacking. As we show in Section 5 below, the same property also explains why discriminatory disclosures may dominate non-discriminatory ones when the primitive heterogeneity in structural beliefs does not minimize the ex-ante probability of regime change.

In the Supplementary Material, we show that the result in Theorem 1 extends to a fairly general class of economies in which (a) agents' prior beliefs need not be consistent with a common prior, nor be generated by signals drawn independently conditionally on θ , (b) the number of agents is arbitrary (in particular, finitely many agents), (c) agents' may have a level-K degree of sophistication, (d) payoffs may be heterogenous across agents, and (d) the designer need not be able to identify perfectly the "state" (i.e., she may possess imperfect information about the fundamentals θ and/or the agents' beliefs). The key assumptions are (i) the supermodularity of the agents' payoffs, (ii) the measurability of the regime outcome under MARP in the designer's information, and (iii) the invariance of the designer's payoff in the size of attack in the event of regime change.

Finally, it is worth stressing that, although the result in Theorem 1 might be evocative of the *Revelation Principle* (RP henceforth), it is, in fact, fundamentally different. First, the RP does not hold when the solution concept is rationalizability. Second, even if the solution concept were Bayes-Nash Equilibrium, the RP does not hold when the performance of different mechanisms is evaluated on the basis of the most adversarial continuation equilibrium (see Example 1 below for an illustration). Third, the RP says that there is no loss of generality in restricting attention to disclosures that take the form of action recommendations; it does not imply that it is without loss of generality to recommended the same action to all agents. Lastly, the RP does not depend on the payoff structure, whereas the result in Theorem 1 above hinges on the game being supermodular, and on the existence of an aggregate outcome.

As anticipated in the Introduction, the value of Theorem 1 when it comes to banks' stress tests is in identifying the right form of transparency. Optimal tests should not be expected to generate conformism in investors' beliefs about banks' balance sheets (in fact, such conformism can

be detrimental to the possibility of minimizing the risk of defaults). However, there is no value in inducing uncertainty about the market response to the tests or in the fate of the banks.

4 Public Disclosures

We now specialize the analysis to environments where the policy maker is constrained to disclose the same information to all market participants. Formally, public disclosures are non-discriminatory policies $\pi : \Theta \rightarrow \Delta(M(\mathcal{S}))$ such that, for all $\theta \in \Theta$, all $m \in \text{supp}[\pi(\theta)]$, there exists a public signal $s \in \mathcal{S}$ such that $m_i = s$, all $i \in [0, 1]$.

4.1 Simple Pass/Fail Tests

The next theorem identifies conditions under which optimal non-discriminatory policies take the form of simple Pass/Fail tests.

Theorem 2. *Assume $p(x|\theta)$ is log-supermodular. Given any non-discriminatory policy Γ , there exists a binary non-discriminatory policy $\Gamma^* = \{\mathcal{S}^*, \pi^*\}$, $\mathcal{S}^* = \{0, 1\}$, yielding the policy maker a payoff weakly higher than Γ and such that, under MARP associated with Γ^* , when signal $s = 0$ is disclosed, all agents refrain from attacking, whereas when signal $s = 1$ is disclosed, they all attack, irrespective of their exogenous private information.*

The proof in the Appendix is in three steps. First, using arguments similar to those establishing Theorem 1, it shows that, starting from Γ , one can construct another policy $\hat{\Gamma}$ that, in addition to the signals disclosed by Γ , it discloses the regime outcome $r(\theta, s; a^\Gamma) \in \{0, 1\}$ that would have prevailed at (θ, s) under MARP consistent with Γ . The new policy $\hat{\Gamma}$ satisfies the perfect coordination property and weakly improves upon Γ . Step 2 then shows that, when the exogenous signals are drawn from a log-supermodular distribution (equivalently when the signals satisfy the *monotone likelihood ratio property*) then, irrespective of the disclosure policy, MARP takes the form of a cut-off strategy profile according to which, given any s , agents attack if, and only if, their private signals fall below a cutoff ξ_s . Finally, Step 3 builds on Step 2 to show that, starting from $\hat{\Gamma}$, one can construct a new policy Γ^* that, for each message $(s, r(\theta, s; a^\Gamma))$ sent with positive probability under $\hat{\Gamma}$, conceals signal s and only discloses $r(\theta, s; a^\Gamma)$, without changing the agents' behavior.

To see this more precisely, given any arbitrary policy Γ_0 , let $U^{\Gamma_0}(x, (s, 0)|k)$ denote the payoff from not attacking of an agent with exogenous signal x who receives endogenous public information $(s, 0)$ and expects all other agents to follow a cut-off strategy with cut-off k (that is, to attack if, and only if, their private signals fall short of the cut-off k). That the policy $\hat{\Gamma}$ satisfies the perfect coordination property, along with the fact that, when $p(x|\theta)$ is log-supermodular, MARP is in cut-off strategies, implies that, for any public announcement $(s, 0)$, any cutoff k , the payoff from not attacking of any agent whose private signal x coincides with the cutoff k must be strictly positive, that is $U^{\hat{\Gamma}}(k, (s, 0)|k) > 0$. By the law of iterated expectations, the same property is then guaranteed

to be true under the policy Γ^* , for the latter is simply obtained from $\hat{\Gamma}$ by concealing the information s . That is, when the policy Γ^* publicly announces that the result of the test is a pass (formally, $r = 0$), for any cu-off k ,

$$U^{\Gamma^*}(k, 0|k) = \int U^{\hat{\Gamma}}(k, (s, 0)|k) d\Lambda^{\hat{\Gamma}}(s|0, k) > 0,$$

where $\Lambda^{\hat{\Gamma}}(\cdot|0, k)$ is the probability an agent with signal $x = k$, who hears the public announcement that $r = 0$, assigns to the policy $\hat{\Gamma}$ having disclosed the public signal $(s, 0)$ when the signals s are concealed to the agents. This means that, under MARP consistent with the new policy Γ^* , no agent attacks when hearing the public announcement that $r = 0$. The policy maker can thus replace the original policy Γ with the new policy Γ^* and guarantee that, for each fundamental θ , the market response is the same as under the original policy Γ .

The key property that justifies restricting attention to simple pass/fail policies is the log - supermodularity of the signal distribution $p(x|\theta)$. As anticipated in the Introduction, this property, which is formally equivalent to MLRP, means that states with higher θ (equivalently, with stronger fundamentals) are also states in which more agents are expected to have more optimistic beliefs about θ , that is, beliefs that assign higher probability to higher θ (equivalently to stronger fundamentals), in the sense of first-order-stochastic-dominance. As the next example shows, this property of beliefs is essential for the optimality of simple pass/fail tests.²²

Example 1. The fundamental variable θ is drawn from a uniform distribution over $[2/3, 4/3]$. Given θ , each agent $i \in [0, 1]$ receives an exogenous signal $x_i \in \{L, H\}$, drawn independently across agents from a Bernoulli distribution with probability

$$Pr(x_i = L|\theta) = \begin{cases} 2/3 & \text{if } \theta \in [2/3, 5/6) \cup [1, 7/6) \\ 1/3 & \text{otherwise.} \end{cases}$$

Figure 1 illustrates the signals considered in Example 1. Note that the agents' posterior beliefs under the signal structure of Example 1 can be ranked according to First-Order-Stochastic Dominance. Each agent receiving a High signal has posterior beliefs that dominate those of each agent receiving a Low signal. Nonetheless, the ratio $p(H|\theta)/p(L|\theta)$ is not increasing over θ over the entire domain, meaning that $p(x|\theta)$ is not log-supermodular or, equivalently, it does not satisfy the *monotone-likelihood-ratio property*.

Suppose that agents' payoffs are such that $g(\theta, A) = -b(\theta, A) = c$ for all (θ, A) and that $R(\theta, A) = \theta - A$. Under such a specification, attacking is rationalizable if the probability of regime change is at least $1/2$, whereas not attacking is rationalizable if the probability of regime change is no greater than $1/2$. In the absence of additional information, MARP consistent with the information structure described above then features all agents attacking, regardless of their exogenous signals. In fact, if

²²We thank Tommaso Denti for providing us with such example.

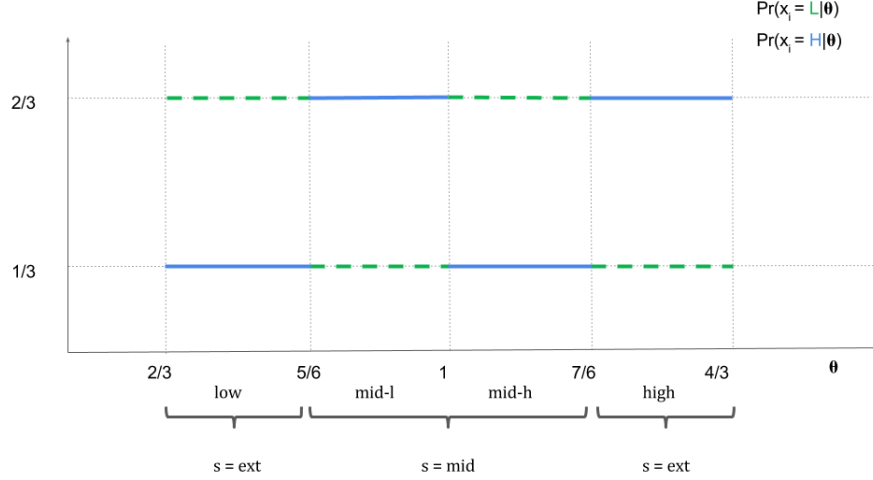


Figure 1: Suboptimality of simple pass/tail tests

all agents attack, then regime change occurs for all $\theta \leq 1$. Given the symmetry of the example, all agents believe that $Pr(\theta \leq 1|x) = 1/2$ irrespective of the realization of x and therefore attacking is rationalizable for all individuals.

Next, observe that the designer can prevent regime change, for any realization of the fundamentals θ , by committing to a non-discriminatory policy $\Gamma^{ex} = (\{mid, ext\}, \pi)$ that publicly announces whether the fundamentals are extreme ($s = ext$) or intermediate ($s = mid$), with

$$\pi(mid|\theta) = \begin{cases} 1 & \text{if } \theta \in [5/6, 7/6] \\ 0 & \text{if otherwise.} \end{cases}$$

See again Figure 1 for a graphical representation of such a policy. To see that, under such a policy, all agents refrain from attacking, no matter θ , consider first the case where the fundamentals are extreme, i.e., $\theta \in [2/3, 5/6] \cup [7/6, 4/3]$. Given the public announcement that $s = ext$, all agents with exogenous signal $x = H$ then find it dominant to refrain from attacking (this is because, even if all other agents were to attack, the probability that each agent observing signal H would assign to regime change would be equal to $Pr(\theta \leq 1|H, ext) = 1/3 < 1/2$ making it optimal for the individual not to attack). As a consequence of such a property, each agent receiving an exogenous signal equal to L would then find it *iteratively* dominant not to attack (this is because, for any $\theta \in [2/3, 5/6]$, even if all agents receiving signals equal to L were to attack, the aggregate size of the attack would be equal to $Pr(L|\theta) = 2/3 < \theta$, implying that the attack would fail, while, for $\theta \in [7/6, 4/3]$, an attack of any size would fail).

Next, consider the case in which the fundamentals are intermediate, i.e., $\theta \in (5/6, 7/6)$. In this case, all agents with signal equal to L assign probability $2/3$ to $\theta \geq 1$ and hence find it dominant not to attack. Because, for any $\theta \in (5/6, 1)$, $1/3$ of the agents receive an L signal, the maximal size of an attack that each agent with signal equal to H can expect at any $\theta \in (5/6, 1)$ is thus equal

to $Pr(H|\theta) = 2/3 < \theta$, implying that even if all agents with signal equal to H were to attack, the attack would fail, thus making it iteratively dominant for the agents receiving the H signal not to attack.

As a result, under Γ^{ex} , at any θ , the unique rationalizable profile features all agents refraining from attacking, regardless of their private information.

Because, under Γ^{ex} , no agent attacks, irrespective of the fundamental θ and of the public signal $s \in \{ext, mid\}$, one may then conjecture that the policy maker could simply recommend to the agents to abstain from attacking and refrain from disclosing whether the fundamentals are intermediate or extreme. In the contest of stress tests, this would amount to a public announcement that the financial institution passed the test, without any further information disclosed to the market. However, when the market is expected to play according to MARP, such a simple announcement would not avoid regime change, for, as discussed above, when the policy does not change the agents' beliefs, the most aggressive rationalizable profile features all agents attacking.

The example thus illustrates both the failure of the Revelation Principle (when the market is expected to play according to MARP, it is with loss of generality to confine attention to policies that take the form of action recommendations), as well as the suboptimality of simple pass/fail tests, when beliefs and fundamentals do not co-move according to the monotone likelihood ratio property.

4.2 Monotone Tests

We now turn to the optimality of policies that fail with certainty institutions with weak fundamentals and pass with certainty those with strong fundamentals. For any $(\theta, x) \in \mathbb{R}^2$, let $B(\theta, x) \equiv b(\theta, P(x|\theta))$ denote the agents' payoff from refraining from attacking when regime change occurs, the fundamentals are equal to θ , and the aggregate size of attack is $A(\theta, x) = P(x|\theta)$.

Condition 1. *The following properties hold:*

1. *The payoff differential $\Delta^P(\theta) \equiv W(\theta, 0) - L(\theta)$ is nondecreasing in θ ;*
2. *The functions $p(x|\theta)$ and $B(\theta, x)$ are log-supermodular;*
3. *For any x , the function $Y(\theta; x) \equiv \Delta^P(\theta)/[p(x|\theta)B(\theta, x)]$ is strictly increasing in θ over $[\underline{\theta}, \hat{\theta}(x)]$, where $\hat{\theta}(x)$ is the regime threshold when agents follow cut-off strategies with cut-off x (i.e., $\hat{\theta}(x)$ solves $R(\hat{\theta}(x), P(x|\hat{\theta}(x))) = 0$).*

Theorem 3. *Suppose Condition (1) holds. Given any non-discriminatory policy Γ , there exists a deterministic monotone non-discriminatory policy $\Gamma^* = (\{0, 1\}, \pi^*)$ satisfying the perfect-coordination property that yields the policy maker a payoff weakly higher than Γ . The policy $\Gamma^* = (\{0, 1\}, \pi^*)$ is defined by a threshold $\theta^* \in [\underline{\theta}, \bar{\theta}]$ such that, for all $\theta \leq \theta^*$, $\pi^*(1|\theta) = 1$, whereas, for all $\theta > \theta^*$, $\pi^*(0|\theta) = 1$.*

That $\Delta^P(\theta)$ is nondecreasing means that the net value the policy maker assigns to moving the economy away from regime change towards a situation in which no agent attacks is monotone in the fundamentals. This condition alone implies that, starting from any non-monotone (and possibly stochastic) policy, there exists a deterministic monotone policy that yields the policy maker a higher payoff. In games with a single receiver sharing the same prior as the policy maker, such condition suffices to guarantee the optimality of threshold policies. This is not the case with multiple receivers with heterogeneous private information. The extra conditions in the theorem guarantee the possibility of constructing perturbations of the original policy by swapping the probability of saving the regime from low to high states while also preserving the agents' incentives not to attack when recommended to do so (equivalently, the uniqueness of the rationalizable profile in the continuation game that follows the announcement the institutions under examination passed the test). As the next example shows, the above condition is quite stringent, in the sense that it need not be satisfied under standard assumptions in the *global games* literature.

Example 2. Suppose (a) the policy maker's payoff is equal to L in case of regime change and W in case of no regime change, independently of (θ, A) , (b) regime change occurs if, and only if, $A \geq \theta$, (c) agents' exogenous signals are given by $x_i = \theta + \sigma \epsilon_i$ where $\sigma \in \mathbb{R}_+$ and where each ϵ_i is drawn from a standard Normal distribution, independently across agents and independently from θ , (d) the agents' payoff are equal to $g = 1 - c$ in case of no regime change and $b = -c$ in case of regime change, with $c > 1/2$. Then, for σ small, Condition (1) is violated and the optimal non-discriminatory policy is necessarily non-monotone.

Note that the property of Condition (1) which is violated under the specification in Example 2 is the third one. The function $Y(\theta; x) \equiv (W - L)/[p(x|\theta)c]$ is not monotone in θ when x is drawn from a Normal distribution. As we show in the Supplementary Material, when σ is small, the designer is then strictly better off by switching from a monotone policy to a non-monotone one that passes institutions with fundamentals either above a critical value θ'' or between $\underline{\theta}$ and θ' and fails the others. Importantly, as the example reveals, the optimality of non-monotone policies does not stem from non-monotonicities in the policy maker's or the agents' payoffs, or from "weird" properties of the agents' beliefs. It originates in the fact that the agents (the receivers in the persuasion game) play according to MARP in the continuation game under asymmetric information. By letting some institutions with intermediate fundamentals fail, the policy maker can save some institutions with lower fundamentals. When property 3 in Condition (1) is violated, the value the policy maker assigns to saving institutions with stronger fundamentals relative to the value it assigns to saving institutions with weaker fundamentals is not large enough to compensate for the fact that, when agents possess private information, the "amount" of institutions that can be saved under non-monotone policies is larger. When this is the case, non-monotone policies can improve upon monotone ones.

In the context of stress test design, the assumption that the differential $\Delta^P(\theta)$ is non-decreasing means that the benefit the policy maker assigns to inducing all investors to pledge to a solvent but

illiquid bank increases with the size of the bank's performing loans, or, more generally, with the value and profitability of the bank's assets. Importantly, stronger supermodularity conditions such as those requiring the function W to be supermodular, and/or, the function L to be non-decreasing are not essential to the result. This assumption seems plausible in many cases of interest. Likewise, the assumption that the distribution $p(x|\theta)$ is log-supermodular also seems plausible. Recall that this condition, which corresponds to the familiar monotone likelihood ratio property, is the same condition that guarantees the optimality of simple pass/fail tests, as shown in Theorem 2. The property that the agents' payoff $B(\theta, x) \equiv b(\theta, P(x|\theta))$ under regime change is log-supermodular in (θ, x) also seems plausible. This property captures the idea that the smaller the fraction of investors pledging to the bank (i.e., the higher x is), the larger the effect of an improvement in the bank's balance sheet (captured by a higher θ) on the reduction in the loss that each investor derives from pledging to a defaulting bank. Unfortunately, as the above example illustrates, the above properties do not guarantee that the optimal policy is monotone. One also needs the value the policy maker assigns to saving a troubled bank to grow with the bank's fundamentals sufficiently fast relative to the payoff that each agent pledging to a defaulting bank derives, scaled by the distribution of the agents' private signals. We see no compelling reason why such condition should be expected to hold in applications. As a result, we conclude that simple monotone tests are unlikely to be optimal in most cases of interest.

When Condition (1) does hold, the choice of the optimal policy reduces to the choice of the largest threshold θ^* such that, for all $x \in \mathbb{R}$,

$$\int_{\theta^*}^{\infty} u(\theta, P(x|\theta))p(x|\theta)f(\theta)d\theta > 0.$$

As we show in the Appendix (see the proof of Lemma 1) the above property implies that when the policy fails all institutions with fundamentals below θ^* and passes all institutions with fundamentals above θ^* , in the continuation game that follows the announcement that the bank under scrutiny passed the test, the unique rationalizable profile features all agents refraining from attacking (i.e., pledging to the bank). The above problem, however, does not have a formal solution.²³ Notwithstanding these complications, with abuse, hereafter, we refer to the threshold policy Γ^* with cut-off

$$\theta^* \equiv \inf\{\theta' : \int_{\theta'}^{\infty} u(\theta, P(x|\theta))p(x|\theta)f(\theta)d\theta > 0 \text{ for all } x \in \mathbb{R}\} \quad (2)$$

as to the optimal monotone policy. The reason for the abuse is that the policy maker can guarantee herself a payoff arbitrarily close to the one associated with a monotone pass/fail test with cut-off equal to θ^* .

²³This was first noticed in Goldstein and Huang (2016).

5 Discriminatory Disclosures

We now return to situations in which the policy maker can disclose different information to different market participants. The purpose of this section is to illustrate the benefits of discriminatory disclosures. To maintain the analysis as simple as possible, we assume that Condition (1) holds. Recall that this condition implies that, if the policy maker were to restrict attention to non-discriminatory policies, the optimal one would be a simple monotone pass/fail test failing with certainty all institutions with fundamentals below θ^* and passing with certainty all the others, with θ^* as defined in (2).

We start by discussing the benefits of discriminatory disclosures. We then turn to situations in which the policy maker can engineer any public disclosure, but is constrained to using Gaussian signals when communicating privately with the agents and identify primitive conditions on payoffs under which public disclosures are optimal.

5.1 Benefits of discriminatory disclosures

Perhaps surprisingly, the reason why discriminatory disclosures can improve upon non-discriminatory ones has little to do with the possibility of tailoring the information disclosed to the agents to their prior beliefs. Discriminatory disclosures outperform non-discriminatory ones because, by enhancing the dispersion of posterior beliefs, they make it harder for the agents to coordinate on a successful attack, thus permitting the policy maker to save a larger set of institutions.

To illustrate the point in the simplest possible way, consider an economy in which the agents' prior beliefs are homogenous (formally, this amounts to assuming the exogenous private signals x are completely uninformative). Next notice that, for any $\hat{\theta}$ such that

$$\int u(\theta, 1) dF(\theta | \theta > \hat{\theta}) \leq 0,$$

the most aggressive rationalizable strategy profile following the public announcement that $\theta > \hat{\theta}$ is such that every agent attacks.²⁴ When the environment satisfies the properties of Condition (1), the optimal non-discriminatory policy is a threshold rule with cut-off equal to²⁵

$$\hat{\theta}^* = \inf \{ \hat{\theta} \in \mathbb{R} \text{ s.t. } \int u(\theta, 1) dF(\theta | \theta > \hat{\theta}) > 0 \}. \quad (3)$$

Suppose now the policy maker, instead of announcing whether θ is above or below some threshold $\hat{\theta}$, sends to each individual a private signal of the form $m_i = \theta + \sigma \xi_i$, where $\sigma \in \mathbb{R}_+$ is a scalar, and where the idiosyncratic terms ξ_i are drawn from a smooth distribution over the entire real line (e.g., a standard Normal distribution), independently across agents, and independently from θ . From

²⁴The notation $F(\theta | \theta > \hat{\theta})$ stands for the common posterior obtained from the prior F by conditioning on the event that $\theta > \hat{\theta}$.

²⁵Here we follow the same convention as in Subsection 4.2 and refer to the optimal non-discriminatory policy as to the threshold policy with cut-off $\hat{\theta}^*$.

standard results in the global games literature, we know that, as the private messages become highly precise (formally, as $\sigma \rightarrow 0^+$), in the absence of any public disclosure, under the most aggressive rationalizable profile, each agent attacks if, and only if, his endogenous private signal falls below a threshold $\theta^{MS} \in (\underline{\theta}, \bar{\theta})$ implicitly defined by the unique solution to²⁶

$$\int_0^1 u(\theta^{MS}, l) dl = 0. \quad (4)$$

The threshold θ^{MS} corresponds to the value of the fundamentals at which an agent who knows θ and holds *Laplacian* beliefs with respect to the size of the attack²⁷ is indifferent between attacking and not attacking. Importantly, θ^{MS} is independent of the initial common prior and of the distribution of the noise terms ξ in the agents' signals. The above result thus implies that, with discriminatory disclosures, the policy maker can always guarantee that regime change never occurs for any $\theta > \theta^{MS}$.

Proposition 1. *Assume the agents possess no exogenous private information about the underlying fundamentals. Let $\hat{\theta}^*$ be the threshold in (3) and θ^{MS} be the threshold in (4). Whenever $\theta^{MS} < \hat{\theta}^*$, discriminatory disclosures strictly improve upon non-discriminatory ones.*

The result follows directly from the arguments preceding the proposition. Because $\hat{\theta}^*$ can be arbitrarily close to $\bar{\theta}$ for particular prior distributions, and because θ^{MS} is invariant in the prior distribution from which θ is drawn, the result in Proposition 1 is relevant in many cases of interest.

As anticipated above, the reason why discriminatory disclosures can improve upon non-discriminatory ones is that they permit the policy maker to enhance the dispersion of the agents' first- and higher-order beliefs. A higher dispersion in turn makes it more difficult for the agents to coordinate on a successful attack. Formally speaking, when beliefs are sufficiently dispersed, an agent receiving a private signal indicating that the regime may collapse under a sufficiently large attack may nonetheless refrain from attacking because he is concerned that many other agents may have received extreme signals indicating that the fundamentals are strong enough to survive an attack of any size. In this case, refraining from attacking may become *iteratively dominant* for this individual. The optimality of discriminatory policies thus follows from a “divide-and-conquer” logic reminiscent of the one in the vertical contracting literature (see, e.g., Segal (2003) and the references therein). Importantly, when discriminatory policies outperform non-discriminatory ones, this is not because they mis-coordinate the response by the market (recall that, by virtue of Theorem 1, the optimal policy always satisfies the perfect-coordination property, irrespectively of whether or not it is discriminatory), but because, by enhancing the heterogeneity in structural uncertainty, they make it difficult for market participants to coordinate on attacking when the planner recommends they abstain from doing so.

²⁶See Morris and Shin (2006).

²⁷This means that the agent believes that the proportion of agents attacking is uniformly distributed over $[0, 1]$.

5.2 On the optimality of non-discriminatory policies

We conclude by showing how the optimality of discriminatory policies depends on properties of the agents' payoffs that relate to the type of claims issued by the banks under scrutiny. To gain on tractability, we consider an environment in which the prior distribution F from which θ is drawn is an improper uniform distribution over the entire real line and where the agents' exogenous private signals are given by $x_i = \theta + \sigma_\eta \eta_i$, with $\eta_i \sim \mathcal{N}(0, 1)$.²⁸ Furthermore, to facilitate the aggregation of the agents' exogenous and endogenous signals into unidimensional statistics, we restrict the policy maker to sending signals of the Gaussian form $\tilde{m}_i = \theta + \sigma_\xi \xi_i$, with $\xi_i \sim \mathcal{N}(0, 1)$, when communicating privately with the agents. Note that the policy maker can engineer any public disclosure of her choice. The restriction to Gaussian signals applies only to the information the policy maker discloses *privately* to the agents, over and above the information conveyed by the public test. In each state θ , the endogenous information $m_i = (\tilde{s}, \tilde{m}_i)$ disclosed to each agent i thus comprises a public signal $\tilde{s}$, along with a private signal $\tilde{m}_i$. The quality of the private signals is then conveniently parametrized by the variance $\sigma_\xi^2 > 0$ of the endogenous noise terms.

We also assume the agents' payoff from not attacking is invariant in A , which amounts to assuming that there exist strictly increasing functions $\bar{g}(\theta)$ and $\bar{b}(\theta)$ such that $g(\theta, A) = \bar{g}(\theta)$ and $b(\theta, A) = \bar{b}(\theta)$, all (θ, A) . Finally, we assume here the function R determining the regime outcome takes the familiar linear form $R(\theta, A) = \theta - A$.²⁹

Then observe that the information contained in each pair $(x_i, \tilde{m}_i)$ is the same as the information contained in the sufficient statistics

$$z_i \equiv \frac{\sigma_\xi^2 x_i + \sigma_\eta^2 \tilde{m}_i}{\sigma_\eta^2 + \sigma_\xi^2}, \quad (5)$$

which, given θ , is normally distributed with mean θ and variance $\sigma_z^2 \equiv (\sigma_\eta^2 \sigma_\xi^2) / (\sigma_\eta^2 + \sigma_\xi^2)$. Hence, the policy maker's choice of the discriminatory component of her disclosure policy can be conveniently reduced to the choice of the standard deviation σ_z of the above sufficient statistics, with $\sigma_z \in (0, \sigma_\eta]$.

Arguments analogous to those establishing Lemma 1 in the Appendix then imply that, for any realization $\tilde{s}$ of the endogenous public signal, the most aggressive rationalizable strategy profile a^Γ is characterized by a unique cut-off $\bar{z}(\tilde{s})$ (whose value depends on the distribution from which the public signal is drawn) such that, for all $i \in [0, 1]$, $a_i^\Gamma(x_i, (\tilde{s}, \tilde{m}_i)) = \mathbb{I}\{z_i \leq \bar{z}(\tilde{s})\}$. Moreover, arguments similar to those establishing Theorem 2 above imply that, for any given choice of σ_z^2 , the optimal public announcement is binary with $\tilde{s} \in \{0, 1\}$ — that is, the public test has a pass/fail structure. Finally, from Theorem 1, the optimal policy has the perfect-coordination property which means that, given σ_z^2 , $\bar{z}(0) = -\infty$, and $\bar{z}(1) = +\infty$. That is, all agents attack when $\tilde{s} = 1$, and they all refrain from attacking when $\tilde{s} = 0$.

²⁸The assumption that F is an improper uniform distribution is standard in the global-game literature. It simplifies the formulas below, without any serious effect on the results.

²⁹The results below extend to more general payoff functions, as long as the agents' exogenous signals x are sufficiently precise.

Next, let Φ denotes the cdf of the standard Normal distribution, and define

$$z_{\sigma_z}^*(\theta) \equiv \theta + \sigma_z \Phi^{-1}(\theta), \quad (6)$$

to be the private statistics threshold such that, when all agents attack for $z_i < z_{\sigma_z}^*(\theta)$ and refrain from attacking for $z_i > z_{\sigma_z}^*(\theta)$, regime change occurs when the fundamentals fall below θ and does not occur when they are above θ .³⁰

For any $(\theta_0, \hat{\theta}, \sigma_z)$, let $\psi(\theta_0, \hat{\theta}, \sigma_z)$ denote the payoff from not attacking of an agent with private statistics $z_{\sigma_z}^*(\theta_0)$, when regime change occurs for all $\theta \leq \theta_0 \in [0, 1]$, public information reveals that $\theta \geq \hat{\theta}$, and the total precision of private information is σ_z^{-2} . Then let

$$\theta_{\sigma_z}^{inf} \equiv \inf \left\{ \hat{\theta} : \psi(\theta_0, \hat{\theta}, \sigma_z) > 0 \text{ all } \theta_0 \right\}.$$

Note that, for any $\hat{\theta} > \theta_{\sigma_z}^{inf}$, under the most aggressive rationalizable strategy profile, no agent attacks after the public signal reveals that $\theta \geq \hat{\theta}$. Hereafter, we assume that all agents refrain from attacking also when public disclosures reveal that $\theta \geq \theta_{\sigma_z}^{inf}$. This simplifies the exposition below by permitting us to talk about the “optimal policy.” As discussed above, the latter does not formally exist when agents are expected to play according to the most aggressive rationalizable profile. However, because the policy maker can always guarantee that, no matter the selection of the rationalizable strategy profile, no agent attacks for any $\theta > \theta_{\sigma_z}^{inf}$, we find the abuse justified.

Proposition 2. *Suppose the policy maker is constrained to using Gaussian signals when communicating privately with the market. Let*

$$\sigma_z^* \equiv \arg \min_{\sigma_z \in (0, \sigma_\eta]} \theta_{\sigma_z}^{inf}.$$

The optimal disclosure policy has the following structure. The policy maker publicly announces whether $\theta < \theta_{\sigma_z^}^{inf}$, or whether $\theta \geq \theta_{\sigma_z^*}^{inf}$. In addition, when $\theta \geq \theta_{\sigma_z^*}^{inf}$, the policy maker sends a Gaussian private signal to each agent of precision $\sigma_\xi^{-2} = [\sigma_\eta^2 - (\sigma_z^*)^2]/(\sigma_z^*)^2 \sigma_\eta^2$.*

The result follows from the arguments preceding the proposition – note that the precision of the endogenous private information σ_ξ^{-2} in the proposition is the one that, together with the precision of the exogenous signals σ_η^{-2} yields a total precision σ_z^{-2} for the sufficient statistics z_i that minimizes the threshold $\theta_{\sigma_z}^{inf}$ defining the regime outcome.

Equipped with the result in Proposition 2, we can then identify primitive conditions under which the optimal policy is non-discriminatory. By virtue of Proposition 2, discriminatory disclosures dominate non-discriminatory ones if, and only if, $\sigma_z^* < \sigma_\eta$ (equivalently, if, and only if, there exists $\sigma_z < \sigma_\eta$ such that $\theta_{\sigma_z}^{inf} < \theta_{\sigma_\eta}^{inf}$). For any precision σ_z^{-2} of the agents’ private statistics, let $\theta_{\sigma_z}^\#$ denote the unique solution to the equation $\psi(\theta_{\sigma_z}^\#, \theta_{\sigma_z}^{inf}, \sigma_z) = 0$. Note that, under MARP, $\theta_{\sigma_z}^\#$ identifies

³⁰Given that $R(\theta, A) = \theta - A$, $z_{\sigma_z}^*(\theta)$ is implicitly defined by the equation $\Phi\left(\frac{z^* - \theta}{\sigma_z}\right) = \theta$.

the fundamental threshold below which regime change occurs when the total precision of the agents' private information is σ_z^{-2} , and the endogenous disclosure of public information reveals that $\theta \geq \theta_{\sigma_z}^{inf}$. Let³¹

$$D(\theta, \theta_{\sigma_z}^{\#}) \equiv \begin{cases} \bar{b}'(\theta) & \text{if } \theta < \theta_{\sigma_z}^{\#} \\ \bar{g}'(\theta) & \text{if } \theta \geq \theta_{\sigma_z}^{\#} \end{cases}$$

Proposition 3. *Suppose that, for any $\sigma_z \in [0, \sigma_\eta]$,*

$$\mathbb{E}[D(\theta, \theta_{\sigma_z}^{\#})(\theta - \theta_{\sigma_z}^{\#}) | z^*(\theta_{\sigma_z}^{\#}), \theta \geq \theta_{\sigma_z}^{inf}; \sigma_z] \geq 0. \quad (7)$$

Then the optimal policy is non-discriminatory.

The condition in Proposition 3 is a measure of the sensitivity of the marginal agent's net payoff from not attacking to the underlying fundamentals.³² To see this, note that the condition is equivalent to

$$\frac{\mathbb{E}[\bar{g}'(\theta)(\theta - \theta_{\sigma_z}^{\#}) | z^*(\theta_{\sigma_z}^{\#}), \theta \geq \theta_{\sigma_z}^{\#}; \sigma_z]}{\mathbb{E}[\bar{g}(\theta) | z^*(\theta_{\sigma_z}^{\#}), \theta \geq \theta_{\sigma_z}^{\#}; \sigma_z]} \geq \frac{\mathbb{E}[\bar{b}'(\theta)(\theta - \theta_{\sigma_z}^{\#}) | z^*(\theta_{\sigma_z}^{\#}), \theta \in (\theta_{\sigma_z}^{inf}, \theta_{\sigma_z}^{\#}); \sigma_z]}{\mathbb{E}[\bar{b}(\theta) | z^*(\theta_{\sigma_z}^{\#}), \theta \in (\theta_{\sigma_z}^{inf}, \theta_{\sigma_z}^{\#}); \sigma_z]}.$$

The left-hand side is the elasticity of the marginal agent's expected net payoff from refraining from attacking with respect to the underlying fundamentals, in case the status quo is preserved. The right-hand side is the corresponding elasticity in case of regime change.³³

To gather some intuition, consider the case in which, when the regime collapses, the agents' payoff differential between not attacking and attacking is constant in the underlying fundamentals (i.e., $\bar{b}'(\theta) = 0$ for all θ). In this case, the marginal agent faces only *upside risk*. Hence, when the quality of private information decreases (which amounts to a mean-preserving increase in risk), the agent's expected net payoff from not attacking increases. Starting from any policy that discloses private information to the agents (i.e., for which $\sigma_z < \sigma_\eta$), the policy maker can then do better by reducing the precision of the agents' private information. In this case, the optimal policy is non-discriminatory.

The value of Proposition 3 is in indicating how the optimality of enhancing the dispersion of the agents' posterior beliefs relates to the sensitivity of the agents' payoffs to the underlying fundamentals in case of regime change relative to the corresponding sensitivity in the absence of regime change. In turn, such sensitivity typically depends on the type of claims held by the agents. For example, in the context of stress test design, the above condition holds when investors are *equity holders*. In this case, when regime change occurs (i.e., when the bank defaults), their claims are junior (i.e., subordinated)

³¹Here $\bar{b}'$ and $\bar{g}'$ denote the derivatives of the $\bar{b}$ and $\bar{g}$ functions, respectively.

³²The marginal agent is the one with signal $z_{\sigma_z}^*(\theta_{\sigma_z}^{\#})$. See also Iachan and Nenov (2015) for a similar condition in a related class of games of regime change.

³³Observe that, for the marginal agent with signal $z_{\sigma_z}^*(\theta_{\sigma_z}^{\#})$,

$$\begin{aligned} &Pr(\theta \geq \theta_{\sigma_z}^{\#} | z_{\sigma_z}^*(\theta_{\sigma_z}^{\#}), \theta \geq \theta_{\sigma_z}^{inf}; \sigma_z) \mathbb{E}[\bar{g}(\theta) | z_{\sigma_z}^*(\theta_{\sigma_z}^{\#}), \theta \geq \theta_{\sigma_z}^{\#}; \sigma_z] = \\ &Pr(\theta \in (\theta_{\sigma_z}^{inf}, \theta_{\sigma_z}^{\#}) | z_{\sigma_z}^*(\theta_{\sigma_z}^{\#}), \theta \geq \theta_{\sigma_z}^{inf}; \sigma_z) \mathbb{E}[\bar{b}(\theta) | z_{\sigma_z}^*(\theta_{\sigma_z}^{\#}), \theta \in (\theta_{\sigma_z}^{inf}, \theta_{\sigma_z}^{\#}); \sigma_z]. \end{aligned}$$

with respect to those from other stake holders with higher seniority (e.g., bond holders). In case of default, their payoff then amounts to a liquidation value that is typically little sensitive to the exact amount of the bank’s performing loans (the bank’s fundamentals). On the contrary, when regime change does not occur (i.e., when the government succeeds in persuading the bank’s equity holders to stay put), the value of the equity-holders’ claims reflect the bank’s long-term profitability, which is sensitive to the amount of the bank’s performing loans. The result in Proposition 3 thus indicates that discriminatory disclosures are more likely to be beneficial when the banks seeking external investment do so by issuing bonds than when they do so by issuing equity.

6 Conclusions

The results in the present paper indicate that, in a large class of games of regime change, the optimal policy completely removes any strategic uncertainty, while retaining, and, in some cases, enhancing, structural uncertainty (that is, the dispersion of beliefs about the underlying fundamentals). Under the optimal policy, each agent can perfectly predict the actions of any other agent, but not the beliefs that rationalize such actions. They also show that, when the policy maker is restricted to disclosing the same information to all market participants, the optimal policy often takes the form of a simple pass/fail test. However, the optimal policy need not be monotone. Indeed, the conditions under which the optimal (non-discriminatory) policy fails institutions with weak fundamentals and passes those with strong ones are quite stringent. Finally, the results illustrate the benefits of disclosing different information to different market participants, and relates such benefits to the type of claims issued by the banks under scrutiny.

The above results are worth extending in several directions. The analysis in the present paper assumes the policy maker knows how the distribution of market beliefs correlates with the banks’ fundamentals. Such knowledge may come from previous experience with banks of similar fundamentals, polls, data on professional forecasters, the IOWA betting markets, and the like. While this is a natural starting point, there are many environments in which it is more appropriate to assume that the designer lacks information about the joint distribution of the underlying fundamentals and market beliefs. In future work, it would be interesting to investigate the optimal disclosure policy in such situations. The idea is to apply a robust (i.e., *undominated* max-min) approach to the designer’s problem, whereby the designer expects (a) Nature to select the information structure that minimizes the planner’s payoff, and (b) the market to coordinate on the most-aggressive rationalizable strategy profile. The emphasis on the designer’s policy to be undominated is key here. Given any disclosure policy, the policy maker’s payoff always attains its minimum when the fundamentals are common knowledge among market participants. A policy is undominated if there exists no other policy that yields a weakly higher payoff *across all possible market beliefs*. The characterization of the set of undominated policies is highly relevant both from a theoretical standpoint and for the associated policy implications.

The analysis in the present paper is also static. However, many of the applications of interest are intrinsically dynamic, with agents coordinating on multiple attacks over time and learning from past attacks (see the discussion in Angeletos et al. (2007)). In future work, it would be interesting to extend the analysis in this direction. When the fundamentals are partially persistent over time, the optimal policy must specify the timing of information disclosures and how the information at each period depends on the agents' behavior in previous periods. Furthermore, an unsuccessful attack in one period may make the status quo more vulnerable in subsequent periods. There are difficulties in extending the analysis to dynamic environments, but the returns are worth the effort.

Finally, the analysis in this paper is conducted by assuming that the fundamentals can be observed by the information designer at no cost. When θ represents information that is private to the financial institution under scrutiny, such an assumption need not be appropriate. In future work, it would also be interesting to investigate the problem of a designer who must solicit information from the financial institution prior to communicating with the market. This creates an interesting screening+persuasion problem in the spirit of what is examined in the literature on privacy in sequential contacting (e.g., Calzolari and Pavan (2006a) and Calzolari and Pavan (2006b)).

Appendix

Proof of Theorem 1. Given any regular policy $\Gamma = (\mathcal{S}, \pi)$ and any $n \in \mathbb{N}$, let $T_{(n)}^\Gamma$ be the set of strategies surviving n rounds of IDISDS in the continuation game that starts after the policy maker announces the policy Γ , with $T_{(0)}^\Gamma$ denoting the entire set of strategy profiles $a = (a_i(\cdot))_{i \in [0,1]}$, where for any $i \in [0,1]$, $a_i : \mathbb{R} \times \mathcal{S} \rightarrow [0,1]$. Let $a_{(n)}^\Gamma \equiv (a_{(n),i}^\Gamma(\cdot))_{i \in [0,1]} \in T_{(n)}^\Gamma$ denote the profile in $T_{(n)}^\Gamma$ that minimizes the policy maker's ex-ante payoff. Such a profile also minimizes the policy maker's interim payoff, as it will become clear from the arguments below. Hereafter, we refer to this profile as the most aggressive profile surviving n rounds of IDISDS. The profiles $(a_{(n)}^\Gamma)_{n \in \mathbb{N}}$ can be constructed inductively as follows. The profile $a_{(0)}^\Gamma \equiv (a_{(0),i}^\Gamma(\cdot))_{i \in [0,1]}$ prescribes that all agents attack irrespective of their exogenous and endogenous signals; that is, each $a_{(0),i}^\Gamma(\cdot)$ is such that $a_{(0),i}^\Gamma(x_i, m_i) = 1$, for all $(x_i, m_i) \in \mathbb{R} \times \mathcal{S}$.³⁴ With an abuse of notation, let $U_i^\Gamma((x_i, m_i); a)$ denote the payoff that, under Γ , agent i , with exogenous signal x_i and endogenous signal m_i , obtains from not attacking when all other agents follow the strategy in the profile a . The most aggressive strategy profile surviving n rounds of IDISDS is the one specifying, for each agent $i \in [0,1]$, $a_{(n),i}^\Gamma(x_i, m_i) = 0$ if $U_i^\Gamma((x_i, m_i); a_{(n-1)}^\Gamma) > 0$ and $a_{(n),i}^\Gamma(x_i, m_i) = 1$ if $U_i^\Gamma((x_i, m_i); a_{(n-1)}^\Gamma) \leq 0$. The most aggressive rationalizable strategy profile (MARP) consistent with the policy Γ is then the profile $a^\Gamma = (a_i^\Gamma(\cdot))_{i \in [0,1]} \in T_\infty^\Gamma$ given by

$$a_i^\Gamma(\cdot) = \lim_{n \rightarrow \infty} a_{(n),i}^\Gamma(\cdot), \text{ all } i \in [0,1].$$

³⁴Note that, to ease the notation, we let each individual strategy prescribe an action for all $(x_i, m_i) \in \mathbb{R} \times \mathcal{S}$, including those that need not be consistent with the policy Γ .

Next, consider the policy $\Gamma^+ = (\mathcal{S}^+, \pi^+)$, $\mathcal{S}' \equiv \mathcal{S} \times \{0, 1\}$, obtained from the original policy Γ by replacing each message function $m : [0, 1] \rightarrow \mathcal{S}$ in the support of each lottery $\pi(\theta)$ with the message function $m^+ : [0, 1] \rightarrow \mathcal{S}'$ that discloses to each agent $i \in [0, 1]$ the same message m_i disclosed by the original policy Γ , along with the regime outcome $r(\theta, m; a^\Gamma)$ that would have prevailed under Γ at (θ, m) when all agents play according to MARP consistent with the original policy Γ (that is, $m_i^+ = (m_i, r(\theta, m; a^\Gamma))$). Note that the new policy $\Gamma^+ = (\mathcal{S}^+, \pi^+)$ selects, for each θ , the message function m^+ with the same probability as the original policy Γ selects the function m .

Now, let $T_{(n)}^{\Gamma^+}$ denote the set of strategies surviving n rounds of IDISDS under the new policy Γ^+ , and $a_{(n)}^{\Gamma^+} \equiv \left(a_{(n),i}^{\Gamma^+}(\cdot) \right)_{i \in [0,1]} \in T_{(n)}^{\Gamma^+}$ the profile in $T_{(n)}^{\Gamma^+}$ that minimizes the policy maker's ex-ante payoff, with $a_{(0)}^{\Gamma^+} \equiv \left(a_{(0),i}^{\Gamma^+}(\cdot) \right)_{i \in [0,1]}$ prescribing that all agents attack irrespective of their exogenous and endogenous signals.

Step 1. First, we prove that, for each $i \in [0, 1]$,

$$\{(x_i, m_i) : U_i^\Gamma(x_i, m_i; a) > 0 \ \forall a\} \subseteq \{(x_i, m_i) : U_i^{\Gamma^+}(x_i, (m_i, 0); a) > 0 \ \forall a\}.$$

That is, any agent i who, under Γ , finds it dominant not to attack, given the information (x_i, m_i) , also finds it dominant not to attack under Γ^+ when receiving information $(x_i, (m_i, 0))$.

To see this, first use the fact that the game is supermodular to observe that

$$\{(x_i, m_i) : U_i^\Gamma(x_i, m_i; a) > 0 \ \forall a\} = \left\{ (x_i, m_i) : U_i^\Gamma \left(x_i, m_i; a_{(0)}^\Gamma \right) > 0 \right\}.$$

Likewise,

$$\{(x_i, m_i) : U_i^{\Gamma^+}(x_i, (m_i, 0); a) > 0 \ \forall a\} = \left\{ (x_i, m_i) : U_i^{\Gamma^+} \left(x_i, (m_i, 0); a_{(0)}^{\Gamma^+} \right) > 0 \right\}.$$

Recall that, under both $a_{(0)}^\Gamma$ and $a_{(0)}^{\Gamma^+}$, all agents attack regardless of their exogenous and endogenous information and therefore, under both $a_{(0)}^\Gamma$ and $a_{(0)}^{\Gamma^+}$, regime change occurs if, and only if, $\theta \leq \bar{\theta}$.

Now let $\Lambda_i^\Gamma(\theta, \mathbf{x}, m | x_i, m_i)$ denote the beliefs of agent $i \in [0, 1]$ over $\Theta \times \mathbb{R}^{[0,1]} \times \mathcal{S}^{[0,1]}$ when receiving information $(x_i, m_i) \in \mathbb{R} \times \mathcal{S}$ under Γ , and $\Lambda_i^{\Gamma^+}(\theta, \mathbf{x}, m | x_i, (m_i, 0))$ the corresponding beliefs under Γ^+ after receiving information (x_i, m_i^+) , with $m_i^+ = (m_i, 0)$. Bayesian updating implies that, under Γ^+ ,

$$\partial \Lambda_i^{\Gamma^+}(\theta, \mathbf{x}, m | x_i, (m_i, 0)) = \frac{\mathbb{I}_{\{r(\theta, m; a^\Gamma)=0\}}}{\pi_i^\Gamma(0 | x_i, m_i)} \partial \Lambda_i^\Gamma(\theta, \mathbf{x}, m | x_i, m_i), \quad (8)$$

where $\mathbb{I}_{\{r(\theta, m; a^\Gamma)=0\}}$ is the indicator function, taking value 1 if (θ, m) is such that (a) $m \in \text{supp}[\pi(\theta)]$ and (b) $r(\theta, m; a^\Gamma) = 0$, and taking value 0 otherwise, and where

$$\pi_i^\Gamma(0 | x_i, m_i) \equiv \int_{\{(\theta, \mathbf{x}, m) : r(\theta, m; a^\Gamma)=0\}} d\Lambda_i^\Gamma(\theta, \mathbf{x}, m | x_i, m_i)$$

is the total probability assigned, under Γ , by agent i , with information (x_i, m_i) , to the event $\{(\theta, m) : r(\theta, m; a^\Gamma) = 0\}$. The agents' beliefs under the new policy Γ^+ thus correspond to “truncations” of the beliefs under the original policy Γ .

Then take any $i \in [0, 1]$ and $(x_i, m_i) \in \mathbb{R} \times \mathcal{S}$ such that

$$U_i^\Gamma \left((x_i, m_i); a_{(0)}^\Gamma \right) = \int_{(\theta, \mathbf{x}, m)} \left(b(\theta, 1) \mathbb{I}_{\{\theta \leq \bar{\theta}\}} + g(\theta, 1) \mathbb{I}_{\{\theta > \bar{\theta}\}} \right) d\Lambda_i^\Gamma(\theta, \mathbf{x}, m | x_i, m_i) > 0.$$

The aforementioned property of Bayesian updating implies that, under Γ^+ ,

$$\begin{aligned} U_i^{\Gamma^+} \left((x_i, (m_i, 0)); a_{(0)}^{\Gamma^+} \right) &= \frac{1}{\pi_i^\Gamma(0 | x_i, m_i)} \int_{(\theta, \mathbf{x}, m)} \left(b(\theta, 1) \mathbb{I}_{\{\theta \leq \bar{\theta}\}} + g(\theta, 1) \mathbb{I}_{\{\theta > \bar{\theta}\}} \right) \times \\ &\quad \times \mathbb{I}_{\{r(\theta, m; a^\Gamma) = 0\}} d\Lambda_i^\Gamma(\theta, \mathbf{x}, m | x_i, m_i) \\ &> \frac{1}{\pi_i^\Gamma(0 | x_i, m_i)} \int_{(\theta, \mathbf{x}, m)} \left(b(\theta, 1) \mathbb{I}_{\{\theta \leq \bar{\theta}\}} + g(\theta, 1) \mathbb{I}_{\{\theta > \bar{\theta}\}} \right) d\Lambda_i^\Gamma(\theta, \mathbf{x}, m | x_i, m_i) \\ &= \frac{1}{\pi_i^\Gamma(0 | x_i, m_i)} U_i^\Gamma \left((x_i, m_i); a_{(0)}^\Gamma \right) > 0, \end{aligned}$$

where the first equality follows from the truncation property of Bayesian updating, the first inequality from the fact that, for all (θ, m) such that (a) $m \in \text{supp}[\pi(\theta)]$ and (b) $r(\theta, m; a^\Gamma) = 1$, $b(\theta, 1) \mathbb{I}_{\{\theta \leq \bar{\theta}\}} + g(\theta, 1) \mathbb{I}_{\{\theta > \bar{\theta}\}} = b(\theta, 1) < 0$, the second equality follows from the definition of $U_i^\Gamma \left((x_i, m_i); a_{(0)}^\Gamma \right)$, and the second inequality from the assumption that $U_i^\Gamma \left((x_i, m_i); a_{(0)}^\Gamma \right) > 0$.

This means that any agent for whom not attacking is dominant under Γ , continues to find it dominant not to attack after receiving information $(x_i, (m_i, 0))$ under Γ^+ .

Step 2. Next, take any $n > 1$. Assume that, for any $1 \leq k \leq n - 1$, any $i \in [0, 1]$,

$$\left\{ (x_i, m_i) : U_i^\Gamma(x_i, m_i; a) > 0 \quad \forall a \in T_{(k-1)}^\Gamma \right\} \subseteq \left\{ (x_i, m_i) : U_i^{\Gamma^+}(x_i, (m_i, 0); a) > 0, \quad \forall a \in T_{(k-1)}^{\Gamma^+} \right\}. \quad (9)$$

Recall that this means that any agent who, given (x_i, m_i) , finds it strictly optimal not to attack when his opponents play any strategy surviving k rounds of IDISDS under Γ continues to find it optimal not to attack when expecting his opponents to play any strategy surviving k rounds of IDISDS under Γ^+ , and, in addition to (x_i, m_i) , he is also informed that (θ, m) is such that $r(\theta, m; a^\Gamma) = 0$. Below we show that that the same property extends to strategies surviving n rounds of IDISDS. That is,

$$\left\{ (x_i, m_i) : U_i^\Gamma(x_i, m_i; a) > 0 \quad \forall a \in T_{(n-1)}^\Gamma \right\} \subseteq \left\{ (x_i, m_i) : U_i^{\Gamma^+}(x_i, (m_i, 0); a) > 0, \quad \forall a \in T_{(n-1)}^{\Gamma^+} \right\}. \quad (10)$$

To see this, use again the fact that the game is supermodular, to observe that

$$\left\{ (x_i, m_i) : U_i^\Gamma(x_i, m_i; a) > 0 \quad \forall a \in T_{(n-1)}^\Gamma \right\} = \left\{ (x_i, m_i) : U_i^\Gamma \left((x_i, m_i); a_{(n-1)}^\Gamma \right) > 0 \right\}$$

and, likewise,

$$\left\{ (x_i, m_i) : U_i^{\Gamma^+}(x_i, (m_i, 0); a) > 0, \quad \forall a \in T_{(n-1)}^{\Gamma^+} \right\} = \left\{ (x_i, m_i) : U_i^{\Gamma^+} \left((x_i, m_i); a_{(n-1)}^{\Gamma^+} \right) > 0 \right\},$$

where recall that $a_{(n-1)}^\Gamma$ (alternatively, $a_{(n-1)}^{\Gamma^+}$) is the most aggressive profile surviving $n - 1$ rounds of IDISDS under Γ (alternatively, Γ^+).

Now let $A(\theta, \mathbf{x}, m; a_{(n-1)}^\Gamma)$ and $r(\theta, \mathbf{x}, m; a_{(n-1)}^\Gamma)$ denote, respectively, the aggregate size of attack and the regime outcome that prevail at $(\theta, \mathbf{x}, m)$ when agents play according to $a_{(n-1)}^\Gamma$. Then take any $i \in [0, 1]$ and any $(x_i, m_i) \in \mathbb{R} \times \mathcal{S}$ such that

$$U_i^\Gamma((x_i, m_i); a_{(n-1)}^\Gamma) = \int_{(\theta, \mathbf{x}, m)} u\left(\theta, A(\theta, \mathbf{x}, m; a_{(n-1)}^\Gamma)\right) d\Lambda_i^\Gamma(\theta, \mathbf{x}, m | x_i, m_i) > 0.$$

Because $a_{(n-1)}^\Gamma$ is necessarily more aggressive than a^Γ , for all $(\theta, \mathbf{x}, m)$ such that $\mathbf{x} \in \mathbf{X}(\theta)$ and $m \in \text{supp}[\pi(\theta)]$, $r(\theta, m; a^\Gamma) = r(\theta, \mathbf{x}, m; a^\Gamma) = 1 \Rightarrow r(\theta, \mathbf{x}, m; a_{(n-1)}^\Gamma) = 1$. This means that

$$\begin{aligned} \int_{(\theta, \mathbf{x}, m)} u\left(\theta, A(\theta, \mathbf{x}, m; a_{(n-1)}^\Gamma)\right) \mathbb{I}_{\{r(\theta, m; a^\Gamma)=1\}} d\Lambda_i^\Gamma(\theta, \mathbf{x}, m | x_i, m_i) &= \\ \int_{(\theta, \mathbf{x}, m)} b\left(\theta, A(\theta, \mathbf{x}, m; a_{(n-1)}^\Gamma)\right) \mathbb{I}_{\{r(\theta, m; a^\Gamma)=1\}} d\Lambda_i^\Gamma(\theta, \mathbf{x}, m | x_i, m_i) &< 0. \end{aligned} \quad (11)$$

This observation, together with the truncation property in (8), imply that

$$\begin{aligned} U_i^{\Gamma^+}(x_i, (m_i, 0); a_{(n-1)}^\Gamma) &= \int_{(\theta, \mathbf{x}, m)} u\left(\theta, A(\theta, \mathbf{x}, m; a_{(n-1)}^\Gamma)\right) d\Lambda_i^{\Gamma^+}(\theta, \mathbf{x}, m | x_i, m_i) \\ &= \frac{1}{\pi_i^\Gamma(0 | x_i, m_i)} \int_{(\theta, \mathbf{x}, m)} u\left(\theta, A(\theta, \mathbf{x}, m; a_{(n-1)}^\Gamma)\right) \times \\ &\quad \times \mathbb{I}_{\{r(\theta, m; a^\Gamma)=0\}} d\Lambda_i^\Gamma(\theta, \mathbf{x}, m | x_i, m_i) \\ &> \frac{1}{\pi_i^\Gamma(0 | x_i, m_i)} \int_{(\theta, \mathbf{x}, m)} u\left(\theta, A(\theta, \mathbf{x}, m; a_{(n-1)}^\Gamma)\right) d\Lambda_i^\Gamma(\theta, \mathbf{x}, m | x_i, m_i) \\ &= \frac{1}{\pi_i^\Gamma(0 | x_i, m_i)} U_i^\Gamma((x_i, m_i); a_{(n-1)}^\Gamma) > 0, \end{aligned} \quad (12)$$

where the first and third equalities are by definition, the second equality follows from (8), the first inequality follows from (11), and the last inequality is by assumption.

Next, note that $a_{(n-1)}^\Gamma$ and $a_{(n-1)}^{\Gamma^+}$ are such that, for all $i \in [0, 1]$, all $(x_i, m_i) \in \mathbb{R} \times \mathcal{S}$, $a_{(n-1),i}^\Gamma(x_i, m_i)$, $a_{(n-1),i}^{\Gamma^+}(x_i, (m_i, 0)) \in \{0, 1\}$ and

$$\{(x_i, m_i) : a_{(n-1),i}^\Gamma(x_i, m_i) = 0\} = \{(x_i, m_i) : U_i^\Gamma(x_i, m_i; \bar{a}_{(n-2)}^\Gamma) > 0\}$$

and

$$\{(x_i, m_i) : a_{(n-1),i}^{\Gamma^+}(x_i, (m_i, 0)) = 0\} = \{(x_i, m_i) : U_i^{\Gamma^+}(x_i, (m_i, 0); a_{(n-2)}^{\Gamma^+}) > 0\}.$$

The result in Step 1, along with (9), thus imply that $a_{(n-1)}^\Gamma$ and $a_{(n-1)}^{\Gamma^+}$ are thus such that, for all $i \in [0, 1]$, all $(x_i, m_i) \in \mathbb{R} \times \mathcal{S}$,

$$a_{(n-1),i}^\Gamma(x_i, m_i) = 0 \Rightarrow a_{(n-1),i}^{\Gamma^+}(x_i, (m_i, 0)) = 0. \quad (13)$$

Condition (13), along with the fact that the game is supermodular, implies that

$$U_i^{\Gamma^+}(x_i, (m_i, 0); a_{(n-1)}^\Gamma) > 0 \Rightarrow U_i^{\Gamma^+}(x_i, (m_i, 0); a_{(n-1)}^{\Gamma^+}) > 0. \quad (14)$$

Together (12) and (14) imply the property in (10).

Step 3. Equipped with the results in steps 1 and 2 above, we now prove that, for all $\theta \in \Theta$, all $m \in \text{supp}[\pi(\theta)]$ such that $r(\theta, m; a^\Gamma) = 0$, all $\mathbf{x} \in \mathbf{X}(\theta)$, all $i \in [0, 1]$, $a_i^{\Gamma^+}(x_i, (m_i, 0)) \equiv \lim_{n \rightarrow \infty} a_{(n),i}^{\Gamma^+}(x_i, (m_i, 0)) = 0$. This follows directly from the fact that, as shown above,

$$a_i^\Gamma(x_i, m_i) = 0 \Rightarrow a_i^{\Gamma^+}(x_i, (m_i, 0)) = 0, \quad (15)$$

The announcement that (θ, m) is such that, no matter $\mathbf{x} \in \mathbf{X}(\theta)$, $r(\theta, \mathbf{x}, m; a^\Gamma) = 0$ thus reveals to each agent that $(\theta, \mathbf{x}, m)$ is such that, when agents play according to a^{Γ^+} , regime change does not occur. Because the payoff from refraining from attacking is strictly positive when the regime survives, any agent i receiving information $(x_i, m_i, 0)$ thus necessarily refrains from attacking. Under the new signal structure Γ^+ , all agents thus refrain from attacking, regardless of their exogenous and endogenous private signals, when they learn that (θ, m) is such that $r(\theta, m; a^\Gamma) = 0$. That they all attack when they learn that (θ, m) is such that $r(\theta, m; a^\Gamma) = 1$ follows from the fact that such announcements make it common certainty that $\theta \leq \bar{\theta}$.

We conclude that the new policy Γ^+ satisfies the perfect coordination property. That such a policy improves upon the original policy Γ follows from the fact that, for any θ , the probability of regime change under Γ^+ is the same as under Γ , but, in case the status quo survives, the aggregate attack is smaller under Γ^+ than under Γ . Finally, that such a policy is regular follows from the fact that, for any θ , any $m \in \text{supp}[\pi(\theta)]$, the regime outcome at θ under the new policy Γ^+ when the information disclosed to each agent $i \in [0, 1]$ is $m_i^+ = (m_i, r(\theta, m; a^\Gamma))$ where m_i is the message specified by the function m , and all agents play according to MARP, a^{Γ^+} , coincides with the regime outcome under the original policy Γ and hence is the same across all $\mathbf{x} \in \mathbf{X}(\theta)$. The result in the theorem then follows by taking $\Gamma^* = \Gamma^+$. Q.E.D.

Proof of Theorem 2. The proof is in 3 steps. Step 1 shows that, starting from any non-discriminatory policy Γ , one can construct another non-discriminatory policy $\hat{\Gamma}$ that satisfies the perfect coordination property and yields the policy maker a payoff weakly higher than Γ . Step 2 then shows that, when the signals x are drawn from a log-supermodular distribution, irrespective of the disclosure policy, MARP takes the form of a cut-off strategy profile. Finally, Step 3 uses the property in Step 2 to show that, starting from the policy $\hat{\Gamma}$ constructed in Step 2, one can drop any signal other than the expected regime outcome without changing the agents' behavior.

Step 1. Let $\hat{\Gamma} = \{\hat{\mathcal{S}}, \hat{\pi}\}$, $\hat{\mathcal{S}} = \mathcal{S} \times \{0, 1\}$, be the non-discriminatory policy that, for any θ , discloses the same information as the original policy Γ , along with the regime outcome $r(\theta, s; a^\Gamma) \in \{0, 1\}$ that would have prevailed at $(\theta, s) \in \mathbb{R} \times \mathcal{S}$ under MARP consistent with the original policy Γ . Formally, the new policy $\hat{\Gamma}$ is such that for all $\theta \in \Theta$, all $s \in \mathcal{S}$ and any $r \in \{0, 1\}$, $\hat{\pi}((s, r)|\theta) = \pi(s|\theta)\mathbb{I}\{r = r(\theta, s; a^\Gamma)\}$. Under MARP consistent with $\hat{\Gamma}$, irrespective of the exogenous signal x_i , no agent i attacks after the policy publicly announces $(s, 0)$, and necessarily attacks after the policy publicly announces $(s, 1)$. The proof for this property follows from the same steps establishing Theorem 1,

but adapted to the fact that the policy Γ is non-discriminatory. Clearly, the policy $\hat{\Gamma}$ so constructed satisfies the perfect-coordination property.

Step 2. Next we show that, no matter Γ , when signals are drawn from a log-supermodular $p(x|\theta)$, MARP consistent with Γ is in cut-off strategies. To see this, fix an arbitrary policy $\Gamma = (\mathcal{S}, \pi)$ and, for any pair $(x, s) \in \mathbb{R} \times \mathcal{S}$, let $\Lambda^\Gamma(\theta|x, s)$ represent the endogenous posterior beliefs about θ of each agent receiving exogenous information x and endogenous public information s . Next, let

$$U^\Gamma(x, s|k) = \int_{-\infty}^{\infty} u(\theta, P(k|\theta)) d\Lambda^\Gamma(\theta|x, s),$$

denote the payoff from not attacking of an agent with exogenous private signal x observing an endogenous public signal s , when the rest of the agents follow a cut-off strategy with cut-off k (that is, they attack if, and only if, their private signal falls short of k). We then have the following result:

Lemma 1. Suppose $p(x|\theta)$ is log-supermodular. Given any policy $\Gamma = (\mathcal{S}, \pi)$, MARP consistent with Γ is given by the strategy profile $a^\Gamma \equiv (a_i^\Gamma)_{i \in [0,1]}$ such that, for any $s \in \mathcal{S}$, $x \in \mathbb{R}$, $i \in [0, 1]$, $a_i^\Gamma(x, s) = \mathbb{I}\{x \leq \xi^{\Gamma;s}\}$ with $\xi^{\Gamma;s} \equiv \sup\{x : U^\Gamma(x, s|x) \leq 0\}$, all $s \in \mathcal{S}$. Moreover, the strategy profile a^Γ is a BNE of the continuation game that starts with the announcement of the policy Γ .

Proof of Lemma 1. We first establish the following property (the proof is standard and relegated to the Supplementary Material):

Property 1. Assume the function $g : \mathbb{R}^2 \rightarrow \mathbb{R}_+$ is log-supermodular in (x, θ) and the real-valued function $h : \mathbb{R} \rightarrow \mathbb{R}$ crosses 0 only once from below at $\theta = \theta_0$. Choose any (Lebesgue) measurable subset $\Omega \subseteq \mathbb{R}$ and, for any $x \in \mathbb{R}$, let $\Psi(x; \Omega) \equiv \int_{\Omega} h(\theta)g(x, \theta)d\theta$. Suppose there exists $x^ \in \mathbb{R}$ such that $\Psi(x^*; \Omega) = 0$. Then, necessarily $\Psi(x; \Omega) \geq 0$ for all $x > x^*$, and $\Psi(x; \Omega) \leq 0$ for all $x < x^*$, with both inequalities strict if $\Omega \neq \{\theta : h(\theta) = 0\}$, Ω has strict positive Lebesgue measure, and g is strictly positive and strictly supermodular.*

Now fix the policy $\Gamma = (\mathcal{S}, \pi)$. For any $s \in \mathcal{S}$, then let $\xi_{(1)}^{\Gamma;s} \equiv \sup\{x : U^\Gamma(x, s|\infty) \leq 0\}$. Given the public signal s , it is dominant for any agent with private signal x exceeding $\xi_{(1)}^{\Gamma;s}$ not to attack.

Next, let $T_{(1)}^\Gamma$ denote the set of strategy profiles that survive the first round of IDISDS and let $a_{(1)}^\Gamma \equiv \left(a_{(1),i}^\Gamma\right)_{i \in [0,1]}$ denote the most aggressive profile in $T_{(1)}^\Gamma$, that is, the profile in $T_{(1)}^\Gamma$ that minimizes the policy maker's ex-ante payoff. Then observe that such a profile is given by

$$a_{(1),i}^\Gamma(x, s) = \mathbb{I}\{x \leq \xi_{(1)}^{\Gamma;s}\}, \text{ all } (x, s) \in \mathbb{R} \times \mathcal{S}, \text{ all } i \in [0, 1],$$

and minimizes the policy maker's payoff not just in expectation but for any (θ, s) . This follows from the fact that the expected payoff differential $\int u(\theta, 1)d\Lambda^\Gamma(\theta|x, s)$ between not attacking and attacking crosses 0 only once and from below in x at $x = \xi_{(1)}^{\Gamma;s}$. The single crossing property of $\int u(\theta, 1)d\Lambda^\Gamma(\theta|x, s)$ in turn is a consequence of the fact that $u(\theta, 1)$ crosses 0 only once from below at $\theta = \bar{\theta}$ along with Property 1 above.

Similarly, for any $n > 1$, let $T_{(n)}^\Gamma$ denote the set of strategy profiles that survive n rounds of IDISDS under Γ , and $a_{(n)}^\Gamma \equiv \left(a_{(n),i}^\Gamma\right)_{i \in [0,1]}$ the most aggressive profile in $T_{(n)}^\Gamma$ (that is, the profile in $T_{(n)}^\Gamma$ that minimizes the policy maker's ex-ante payoff). That the continuation game is supermodular, along with the fact that, when agents follow monotone strategies the regime outcome is monotone in θ and the density $p(x|\theta)$ is supermodular implies that, for any $s \in \mathcal{S}$, there exists a unique sequence $(\xi_{(n)}^{\Gamma;s})_n$ such that, for any $n \geq 1$, $a_{(n)}^\Gamma$ is such that

$$a_{(n),i}^\Gamma(x, s) = \mathbb{I}\{x \leq \xi_{(n)}^{\Gamma;s}\}, \text{ all } (x, s) \in \mathbb{R} \times \mathcal{S}, \text{ all } i \in [0, 1],$$

with each $\xi_{(1)}^{\Gamma;s}$ as defined above, and with all other cut-offs $\xi_{(n)}^{\Gamma;s}$, $n > 1$, $s \in \mathcal{S}$, defined inductively by $\xi_{(n)}^{\Gamma;s} \equiv \sup\{x : U^\Gamma(x, s | \xi_{(n-1)}^{\Gamma;s}) \leq 0\}$.

Next, let $T^\Gamma \equiv \cap_{n=1}^\infty T_n^\Gamma$ denote the set of strategy profiles that are *rationalizable* for the agents under the policy Γ . The most aggressive strategy profile in T^Γ is then given by

$$a_i^\Gamma(x, s) \equiv \mathbb{I}\{x \leq \xi^{\Gamma;s}\}, \text{ all } (x, s) \in \mathbb{R} \times \mathcal{S}, \text{ all } i \in [0, 1],$$

where, for any $s \in \mathcal{S}$, $\xi^{\Gamma;s} \equiv \lim_{n \rightarrow \infty} \xi_{(n)}^{\Gamma;s}$. The sequence $(\xi_{(n)}^{\Gamma;s})_n$ is monotone and its limit is given by $\xi^{\Gamma;s} = \sup\{x : U^\Gamma(x, s | x) \leq 0\}$. This establishes the first part of the lemma. That the profile a^Γ is a BNE for the continuation game that starts with the announcement of the policy Γ follows from the fact that, given any $s \in \mathcal{S}$, when all agents follow a cut-off strategy with cutoff $\xi^{\Gamma;s}$, the best response for each agent $i \in [0, 1]$ is to attack for $x_i < \xi^{\Gamma;s}$ and to not attack for $x_i > \xi^{\Gamma;s}$ (he is indifferent for $x_i = \xi^{\Gamma;s}$). Q.E.D.

Step 3. Equipped with the result in Lemma 1, we finally show that, starting from the policy $\hat{\Gamma}$ constructed in Step 1, one can construct a “pass/fail” policy $\Gamma^* = (\{0, 1\}, \pi^*)$ that discloses only the expected regime outcome and that yields the policy maker the same payoff as the policy $\hat{\Gamma}$. The policy $\Gamma^* = (\{0, 1\}, \pi^*)$ is such that, for any θ ,

$$\pi^*(0|\theta) = \sum_{s \in \mathcal{S}} \hat{\pi}(s, 0|\theta) = \sum_{\{s \in \mathcal{S} : r(\theta, s; a^\Gamma) = 0\}} \pi(s|\theta).$$

That is, at each θ , the policy Γ^* recommends not to attack (equivalently, announces a “pass”) with the same probability the original policy Γ would have disclosed signals leading to no regime change under MARP consistent with Γ .

We now show that, under the new policy Γ^* , no agent attacks after the policy publicly announces $s^* = 0$. From Lemma 1, we know that, under the policy Γ^* constructed above, for any cutoff k , any private signal x , the payoff $U^{\Gamma^*}(x, 0|k)$ that any agent with private signal x expects from refraining from attacking after the policy Γ^* publicly announces $s^* = 0$ and all other agents follow a cut-off strategy with cut-off k is equal to

$$U^{\Gamma^*}(x, 0|k) = \int_{\mathcal{S}} U^{\hat{\Gamma}}(x, (s, 0)|k) d\Lambda^{\hat{\Gamma}}(s|x, 0) \quad (16)$$

where $\Lambda^{\hat{\Gamma}}(\cdot|x, 0)$ is the probability distribution over $\mathcal{S}$ obtained by conditioning on the event $(x, 0)$, under the policy $\hat{\Gamma}$. From the argument explained in the first step of the proof, under the policy $\hat{\Gamma}$, for any signal $\hat{s} = (s, 0)$ in the range of $\hat{\pi}$, MARP consistent with $\hat{\Gamma}$ is such that $a_i^{\hat{\Gamma}}(x, (s, 0)) = 0$ all $x \in \mathbb{R}$. This implies that, for all signals $\hat{s} = (s, 0)$ in the range of $\hat{\pi}$, $s \in \mathcal{S}$, all $k \in \mathbb{R}$, $U^{\hat{\Gamma}}(k, (s, 0)|k) > 0$. From (16), we then have that, for all $k \in \mathbb{R}$, $U^{\Gamma^*}(k, 0|k) > 0$. In turn, this implies that, given the policy Γ^* , when the signal $s^* = 0$ is disclosed, under the unique rationalizable profile, no agent attacks, that is, $a_i^{\Gamma^*}(x, 0) = 0$ all x , all $i \in [0, 1]$.

It is also easy to see that, when the policy Γ^* publicly discloses the signal $s^* = 1$, it becomes common certainty among the agents that $\theta \leq \bar{\theta}$. Hence, under MARP consistent with Γ^* , after $s^* = 1$ is disclosed, all agents attack, irrespective of their private signals. The policy Γ^* so constructed thus (a) satisfies the perfect-coordination property, (b) is such that, at any θ , the probability of regime change under Γ^* is the same as under $\hat{\Gamma}$ (and hence the same as under Γ), and (c), in case of no regime change, the size of attack under Γ^* is zero. The statement in the theorem then follows from the above properties. Q.E.D.

Proof of Theorem 3. Without loss of generality, assume the policy $\Gamma = (\mathcal{S}, \pi)$ (a) is a “pass/fail” policy, i.e., $\mathcal{S} = \{0, 1\}$, (b) satisfies the perfect coordination property, and (c) is such that $\pi(0|\theta) = 0$ for all $\theta < \underline{\theta}$, and $\pi(0|\theta) = 1$ for all $\theta > \bar{\theta}$. Note that arguments similar to those establishing Theorem 2 imply that, if Γ does not satisfy these properties, there exists another non-discriminatory policy Γ' that does satisfy these properties and yields the policy maker a payoff weakly higher than Γ . The proof then follows from applying the arguments below to Γ' instead of Γ .³⁵

Clearly, if the original policy $\Gamma = (\{0, 1\}, \pi)$ is such that there exists $\theta^* \in [\underline{\theta}, \bar{\theta}]$ such that $\pi(0|\theta) = 0$ for F -almost all $\theta \leq \theta^*$ and $\pi(0|\theta) = 1$ for F -almost all $\theta \geq \theta^*$, then the threshold policy $\Gamma^* = (\{0, 1\}, \pi^*)$ such that $\pi^*(0|\theta) = 0$ for all $\theta \leq \theta^*$ and $\pi^*(0|\theta) = 1$ for all $\theta > \theta^*$ also satisfies the perfect coordination property and yields the policy maker the same payoff as Γ , in which case the result holds.

Suppose, instead, that Γ is such that there exists no θ^* such that $\pi(0|\theta) = 0$ for F -almost all $\theta \leq \theta^*$ and $\pi(0|\theta) = 1$ for F -almost all $\theta > \theta^*$. We then establish the result by showing that there exist a non-discriminatory threshold policy Γ^* satisfying the perfect coordination property that yields the policy maker’s a payoff at least as high as Γ .

Recall that, for the policy Γ to satisfy the perfect coordination property, it must be that, when the signal $s = 0$ is disclosed, $U^{\Gamma}(x, 0|x) > 0$ all x . Now let $\mathbb{G}$ denote the set of all policies with range $\mathcal{S} = \{0, 1\}$ that, in addition to properties (a)-(c) above, satisfy the additional property that $U^{\Gamma}(x, 0|x) \geq 0$, all x . Let $\tilde{\Gamma} \in \arg \max_{\Gamma' \in \mathbb{G}} \{U^P[\Gamma']\}$ be any policy that maximizes the policy maker’s payoff over the set $\mathbb{G}$.³⁶ Note that such a policy need not satisfy the perfect coordination property

³⁵Recall that, by virtue of Theorem 2, the assumption in property (2) in Condition 1 that $p(x|\theta)$ is log-supermodular implies that the optimal non-discriminatory policy has a “pass/fail” structure.

³⁶The arguments below imply that the set $\arg \max_{\Gamma' \in \mathbb{G}} \{U^P[\Gamma']\}$ is non-empty.

for, under $\tilde{\Gamma}$, there may exist x such that $U^{\tilde{\Gamma}}(x, 0|x) = 0$, implying that, in the continuation game that starts after the policy $\tilde{\Gamma}$ announces $s = 0$, there exists a more aggressive rationalizable profile where each agent attacks if and only if his private signals falls short of x . It is also immediate to see that, because the original policy $\Gamma \in \mathbb{G}$, the policy maker's payoff under the original policy Γ cannot be greater than under $\tilde{\Gamma}$.

Below, we first show that, necessarily, under $\tilde{\Gamma} = (\{0, 1\}, \tilde{\pi})$, there exists θ^* such that $\tilde{\pi}(0|\theta) = 0$ for F -almost all $\theta \leq \theta^*$ and $\tilde{\pi}(0|\theta) = 1$ for F -almost all $\theta > \theta^*$. We then show that the policy maker's payoff under $\tilde{\Gamma}$ can be approximated arbitrarily well by a threshold policy $\Gamma^* \in \mathbb{G}$ that satisfies the perfect coordination property (i.e., such that $U^{\Gamma^*}(x, 0|x) > 0$, all x).

First observe that, under any policy $\Gamma' \in \mathbb{G}$, the function $U^{\Gamma'}(x, 0|x)$ is continuous in x . Now let

$$\mathcal{X} \equiv \left\{x : U^{\tilde{\Gamma}}(x, 0|x) = 0\right\}$$

and then denote by $\bar{x} \equiv \sup \mathcal{X}$ and $\underline{x} \equiv \inf \mathcal{X}$. That $U^{\tilde{\Gamma}}(\cdot, 0|\cdot)$ is continuous implies that $\mathcal{X} \neq \emptyset$ for, otherwise, the policy maker could increase her payoff by increasing $\pi(0|\theta)$ over a set of F -positive probability, thus contradicting the optimality of $\tilde{\Gamma}$.³⁷

Now suppose that, under $\tilde{\Gamma}$, there exists no θ^* such that $\tilde{\pi}(0|\theta) = 0$ for F -almost all $\theta \leq \theta^*$ and $\tilde{\pi}(0|\theta) = 1$ for F -almost all $\theta > \theta^*$. We then show that there exists another policy in $\mathbb{G}$ that strictly improves upon $\tilde{\Gamma}$, thus contradicting the assumption that $\tilde{\Gamma}$ maximizes the policy maker's payoff over $\mathbb{G}$.

Let

$$\theta_0 \equiv \inf\{\theta : \exists \delta > 0 \text{ s.t. } \tilde{\pi}(0|\theta') > 0 \text{ for } F\text{-almost all } \theta' \in [\theta, \theta + \delta)\}$$

and, similarly,

$$\theta_1 = \sup\{\theta : \exists \delta > 0 \text{ s.t. } \tilde{\pi}(0|\theta') < 1 \text{ for } F\text{-almost all } \theta' \in [\theta, \theta + \delta)\}.$$

That, under $\tilde{\Gamma}$, there exists no θ^* with the aforementioned properties implies that $\theta_0 < \theta_1$. Furthermore, $[\theta_0, \theta_1] \subset [\underline{\theta}, \bar{\theta}]$.

Case 1. First, suppose that $\hat{\theta}(\bar{x}) < \theta_1$, where, recall that, for any x , $\hat{\theta}(x)$ is the regime threshold such that, when all agents attack when their signal falls below x and refrain from attacking otherwise, regime change occurs if, and only if, $\theta \leq \hat{\theta}(x)$.

Fix $\epsilon > 0$ small, and let $\delta(\epsilon)$ be implicitly defined by

$$\int_{\theta_0}^{\theta_0 + \epsilon} \tilde{\pi}(0|\theta) dF(\theta) = \int_{\theta_1 - \delta(\epsilon)}^{\theta_1} (1 - \tilde{\pi}(0|\theta)) dF(\theta).$$

³⁷To see this, note that, when $\mathcal{X} = \emptyset$, either $\tilde{\Gamma}$ is a threshold policy with threshold $\theta^* = \underline{\theta}$, or there must exist a set $(\theta', \theta'') \subseteq [\underline{\theta}, \bar{\theta}]$ of F -positive probability over which $\pi(0|\theta) < 1$. Because $W(\theta, 0) - L(\theta) > 0$ over such a set, the policy maker could then increase her payoff by increasing the probability that no attack is recommended over such a set, while guaranteeing that, under MARP consistent with the new policy, agents continue to refrain from attacking when signal $s = 0$ is disclosed.

Consider the policy $\Gamma^\epsilon = (\{0, 1\}, \pi^\epsilon)$ defined by: (a) $\pi^\epsilon(\theta) = \tilde{\pi}(\theta)$ for all $\theta \notin [\theta_0, \theta_0 + \epsilon] \cup [\theta_1 - \delta(\epsilon), \theta_1]$; (b) $\pi^\epsilon(0|\theta) = 0$ for all $\theta \in [\theta_0, \theta_0 + \epsilon]$, and (c) $\pi^\epsilon(0|\theta) = 1$ for all $\theta \in [\theta_1 - \delta(\epsilon), \theta_1]$. It is easy to see that such a policy exists and can be chosen to satisfy $\hat{\theta}(\bar{x}) < \theta_1 - \delta(\epsilon)$. Then observe that the policy Γ^ϵ improves the policy maker's payoff upon $\tilde{\Gamma}$. This follows from the fact that the policy Γ^ϵ preserves the ex-ante probability the status quo survives, along with the fact that the policy maker's payoff differential $\Delta^P(\theta)$ is nondecreasing. To establish the result in the theorem it then suffices to show that the new policy Γ^ϵ belongs in $\mathbb{G}$. To see this, let $B(\theta, x) \equiv b(\theta, P(x|\theta))$ and $G(\theta, x) \equiv g(\theta, P(x|\theta))$ denote the agents' payoffs from not attacking when the fundamentals are θ and all agents follow a cut-off strategy with cut-off x , respectively under regime change and under no regime change. Then let

$$p^{\tilde{\Gamma}}(x, 0) \equiv \int \tilde{\pi}(0|\theta)p(x|\theta)dF(\theta) \text{ and } p^{\Gamma^\epsilon}(x, 0) \equiv \int \pi^\epsilon(0|\theta)p(x|\theta)dF(\theta)$$

denote the density of $(x, 0)$, under $\tilde{\Gamma}$ and Γ^ϵ , respectively. Next, observe that, for any $x \leq \hat{\theta}^{-1}(\theta_0 + \epsilon)$ (i.e., for any x such that $\hat{\theta}(x) < \theta_0 + \epsilon$), $U^{\Gamma^\epsilon}(x, 0|x) > 0$, whereas for any $x \in [\hat{\theta}^{-1}(\theta_0 + \epsilon), \bar{x}]$,

$$\begin{aligned} U^{\Gamma^\epsilon}(x, 0|x)p^{\Gamma^\epsilon}(x, 0) &= \int_{\theta_0 + \epsilon}^{\hat{\theta}(x)} B(\theta, x)\pi^\epsilon(0|\theta)p(x|\theta)dF(\theta) + \int_{\hat{\theta}(x)}^{+\infty} G(\theta, x)\pi^\epsilon(0|\theta)p(x|\theta)dF(\theta) \\ &> \int_{\theta_0}^{\hat{\theta}(x)} B(\theta, x)\tilde{\pi}(0|\theta)p(x|\theta)dF(\theta) + \int_{\hat{\theta}(x)}^{+\infty} G(\theta, x)\pi(0|\theta)p(x|\theta)dF(\theta) \\ &= U^{\tilde{\Gamma}}(x, 0|x)p^{\tilde{\Gamma}}(x, 0) \geq 0. \end{aligned}$$

The first inequality follows from the fact that, for any $\theta \in [\theta_0 + \epsilon, \hat{\theta}(x)]$, $\pi^\epsilon(\theta) = \tilde{\pi}(\theta)$, along with the properties that, for any (θ, x) , $B(\theta, x) < 0 < G(\theta, x)$, and the fact that, for all $\theta \geq \hat{\theta}(x)$, $\pi^\epsilon(\theta) \geq \tilde{\pi}(\theta)$. Hence, for all $x \leq \bar{x}$, $U^{\Gamma^\epsilon}(x, 0|x) > 0$. Now recall that, by definition of $\bar{x}$, $U^{\tilde{\Gamma}}(x, 0|x) > 0$ for all $x > \bar{x}$. That $U^{\Gamma^\epsilon}(x, 0|x)$ is continuous in (x, ϵ) , along with the fact that $U^{\Gamma^\epsilon}(\bar{x}, 0|\bar{x}) > 0$, then also implies that, for ϵ strictly positive, but small enough, $U^{\Gamma^\epsilon}(x, 0|x) \geq 0$ for all $x > \bar{x}$. The new policy Γ^ϵ thus belongs in $\mathbb{G}$, as claimed above.

Case 2. Next, suppose that $\hat{\theta}(\bar{x}) \geq \theta_1$. Fix $\epsilon > 0$ small and let $\delta(\epsilon)$ be implicitly defined by

$$\int_{\theta_0}^{\theta_0 + \epsilon} B(\theta, \bar{x})\tilde{\pi}(0|\theta)p(\bar{x}|\theta)dF(\theta) = \int_{\theta_1 - \delta(\epsilon)}^{\theta_1} B(\theta, \bar{x})(1 - \tilde{\pi}(0|\theta))p(\bar{x}|\theta)dF(\theta). \quad (17)$$

Consider the policy $\Gamma^\epsilon = (\{0, 1\}, \pi^\epsilon)$ defined by the following properties: (a) $\pi^\epsilon(\theta) = \tilde{\pi}(\theta)$ for all $\theta \notin [\theta_0, \theta_0 + \epsilon] \cup [\theta_1 - \delta(\epsilon), \theta_1]$, with $\theta_0 + \epsilon < \theta_1 - \delta(\epsilon)$; (b) $\pi^\epsilon(0|\theta) = 0$ for all $\theta \in [\theta_0, \theta_0 + \epsilon]$; and (c) $\pi^\epsilon(0|\theta) = 1$ for all $\theta \in [\theta_1 - \delta(\epsilon), \theta_1]$.

Note that, for ϵ strictly positive but small, such policy exists. Also note that Condition (17) implies that, under the new policy Γ^ϵ , the payoff from not attacking of an agent with signal $\bar{x}$ who expects all other agents to follow a cut-off strategy with cutoff $\bar{x}$ is the same as under the original policy $\tilde{\Gamma}$. Formally, $U^{\Gamma^\epsilon}(\bar{x}, 0|\bar{x}) = U^{\tilde{\Gamma}}(\bar{x}, 0|\bar{x})$. To see this, recall that, by definition of $\bar{x}$,

$U^{\tilde{\Gamma}}(\bar{x}, 0|\bar{x}) = 0$, or equivalently,

$$\int_{\theta_0}^{\hat{\theta}(\bar{x})} B(\theta, \bar{x}) \tilde{\pi}(0|\theta) p(\bar{x}|\theta) dF(\theta) + \int_{\hat{\theta}(\bar{x})}^{+\infty} G(\theta, \bar{x}) \tilde{\pi}(0|\theta) p(\bar{x}|\theta) dF(\theta) = 0.$$

Condition (17), along with properties (a) and (b) above, imply that

$$\int_{\theta_0+\epsilon}^{\hat{\theta}(\bar{x})} B(\theta, \bar{x}) \pi^\epsilon(0|\theta) p(\bar{x}|\theta) dF(\theta) + \int_{\hat{\theta}(\bar{x})}^{+\infty} G(\theta, \bar{x}) \pi^\epsilon(0|\theta) p(\bar{x}|\theta) dF(\theta) = 0.$$

Because

$$U^{\Gamma^\epsilon}(\bar{x}, 0|\bar{x}) = \frac{\int_{\theta_0+\epsilon}^{\hat{\theta}(\bar{x})} B(\theta, \bar{x}) \pi^\epsilon(0|\theta) p(\bar{x}|\theta) dF(\theta) + \int_{\hat{\theta}(\bar{x})}^{+\infty} G(\theta, \bar{x}) \pi^\epsilon(0|\theta) p(\bar{x}|\theta) dF(\theta)}{p^{\Gamma^\epsilon}(0, \bar{x})}.$$

we conclude that $U^{\Gamma^\epsilon}(\bar{x}, 0|\bar{x}) = 0$, as claimed.

We now establish that, for ϵ sufficiently small, under the new policy Γ^ϵ , $U^{\Gamma^\epsilon}(x, 0|x) > 0$ for any $x \in \mathcal{X}$, with $x \neq \bar{x}$. A necessary and sufficient condition for this to be the case is that, for any such x ,

$$\lim_{\epsilon \rightarrow 0^+} \frac{\partial}{\partial \epsilon} U^{\Gamma^\epsilon}(x, 0|x) > 0. \quad (18)$$

Condition (18) holds if, and only if, for any $x \in \mathcal{X}$, $x \neq \bar{x}$,

$$\lim_{\epsilon \rightarrow 0^+} \frac{\partial}{\partial \epsilon} \{U^{\Gamma^\epsilon}(x, 0|x) p^{\Gamma^\epsilon}(0, x)\} > 0. \quad (19)$$

To see that (19) holds, implicitly differentiate the equation in Condition (17) with respect to ϵ and then take the limit as $\epsilon \rightarrow 0$ to observe that

$$\lim_{\epsilon \rightarrow 0^+} \delta'(\epsilon) = \frac{\tilde{\pi}(0|\theta_0) f(\theta_0) p(\bar{x}|\theta_0) |B(\theta_0, \bar{x})|}{(1 - \tilde{\pi}(0|\theta_1)) f(\theta_1) p(\bar{x}|\theta_1) |B(\theta_1, \bar{x})|}.$$

Now take first any $x \in [\hat{\theta}^{-1}(\theta_1), \bar{x}) \cap \mathcal{X}$ and observe that, for any such x ,

$$\begin{aligned} U^{\Gamma^\epsilon}(x, 0|x) p^{\Gamma^\epsilon}(0, x) &= \int_{\theta_0+\epsilon}^{\theta_1-\delta} B(\theta, x) \tilde{\pi}(0|\theta) p(x|\theta) dF(\theta) + \int_{\theta_1-\delta}^{\theta_1} B(\theta, x) p(x|\theta) dF(\theta) \\ &\quad + \int_{\theta_1}^{\hat{\theta}(x)} B(\theta, x) p(x|\theta) dF(\theta) + \int_{\hat{\theta}(x)}^{+\infty} G(\theta, x) p(x|\theta) dF(\theta). \end{aligned} \quad (20)$$

Now use (20) to observe that, for any $x \in [\hat{\theta}^{-1}(\theta_1), \bar{x}) \cap \mathcal{X}$,

$$\begin{aligned} \lim_{\epsilon \rightarrow 0^+} \frac{\partial}{\partial \epsilon} \{U^{\Gamma^\epsilon}(x, 0|x) p^{\Gamma^\epsilon}(0, x)\} &= -B(\theta_0, x) \tilde{\pi}(0|\theta_0) p(x|\theta_0) f(\theta_0) \\ &\quad + B(\theta_1, x) (1 - \tilde{\pi}(0|\theta_1)) p(x|\theta_1) f(\theta_1) \lim_{\epsilon \rightarrow 0} \delta'(\epsilon) \\ &= \tilde{\pi}(0|\theta_0) f(\theta_0) p(\bar{x}|\theta_0) |B(\theta_0, \bar{x})| \left(\frac{p(x|\theta_0)}{p(\bar{x}|\theta_0)} \frac{|B(\theta_0, x)|}{|B(\theta_0, \bar{x})|} - \frac{p(x|\theta_1)}{p(\bar{x}|\theta_1)} \frac{|B(\theta_1, x)|}{|B(\theta_1, \bar{x})|} \right) > 0, \end{aligned}$$

where the inequality follows from the log-supermodularity of $p(x|\theta)$ and $B(\theta, x)$, as per property (2) in Condition (1).

Finally, consider any $x \in [\underline{x}, \hat{\theta}^{-1}(\theta_1)] \cap \mathcal{X}$. That $U^{\tilde{\Gamma}}(\underline{x}, 0|\underline{x}) = 0$ implies that $\hat{\theta}(\underline{x}) > \theta_0$. Hence, for any $x \in [\underline{x}, \hat{\theta}^{-1}(\theta_1)] \cap \mathcal{X}$, $\hat{\theta}(x) > \theta_0$, which means that

$$\begin{aligned} U^{\Gamma^\epsilon}(x, 0|x)p^{\Gamma^\epsilon}(0, x) &= \mathbb{I}(\hat{\theta}(x) > \theta_0 + \epsilon) \int_{[\theta_0 + \epsilon, \hat{\theta}(x)]} B(\theta, x) \tilde{\pi}(0|\theta) p(x|\theta) dF(\theta) \\ &+ \int_{[\max\{\hat{\theta}(x), \theta_0 + \epsilon\}, \theta_1 - \delta(\epsilon)]} G(\theta, x) \tilde{\pi}(0|\theta) p(x|\theta) dF(\theta) \\ &+ \int_{\theta_1 - \delta(\epsilon)}^{\theta_1} G(\theta, x) p(x|\theta) dF(\theta) + \int_{\theta_1}^{+\infty} G(\theta, x) \tilde{\pi}(0|\theta) p(x|\theta) dF(\theta) \\ &> U^{\tilde{\Gamma}}(x, 0|x)p^{\tilde{\Gamma}}(0, x) > 0, \end{aligned}$$

We conclude that, when ϵ is small, under the new policy Γ^ϵ , for all $x \in \mathcal{X}$, $x \neq \bar{x}$, $U^{\Gamma^\epsilon}(x, 0|x) > 0$. That, under the same policy, $U^{\Gamma^\epsilon}(x, 0|x) \geq 0$ also for $x \notin \mathcal{X}$ follows from the fact that $U^{\Gamma^\epsilon}(x, 0|x)$ is continuous in (x, ϵ) .

From the arguments above, we have that the new policy $\Gamma^\epsilon \in \mathbb{G}$. We now show that, when, in addition, property (3) in Condition (1) holds, the new policy yields the policy maker an expected payoff strictly higher than $\tilde{\Gamma}$. To see this, observe that, for any $\epsilon \geq 0$, the policy maker's payoff under the policy Γ^ϵ is equal to

$$\begin{aligned} U^P[\Gamma^\epsilon] &= \int_{-\infty}^{\theta_0 + \epsilon} L(\theta) dF(\theta) + \int_{\theta_1 - \delta(\epsilon)}^{\theta_1} W(\theta, 0) dF(\theta) \\ &+ \int_{(\theta_0 + \epsilon, \theta_1 - \delta(\epsilon)) \cup (\theta_1, +\infty)} (\tilde{\pi}(0|\theta)W(\theta, 0) + (1 - \tilde{\pi}(0|\theta))L(\theta)) dF(\theta). \end{aligned}$$

Differentiating $U^P[\Gamma^\epsilon]$ with respect to ϵ and taking the limit as $\epsilon \rightarrow 0$, we have that:

$$\begin{aligned} \lim_{\epsilon \rightarrow 0^+} \frac{dU^P[\Gamma^\epsilon]}{d\epsilon} &= f(\theta_1)(1 - \tilde{\pi}(0|\theta_1))\Delta^P(\theta_1) \left(\lim_{\epsilon \rightarrow 0} \delta'(\epsilon) \right) - f(\theta_0)\tilde{\pi}(0|\theta_0)\Delta^P(\theta_0) \\ &= f(\theta_0)\tilde{\pi}(0|\theta_0) \left(\Delta^P(\theta_1) \frac{p(\bar{x}|\theta_0)|B(\theta_0, \bar{x})|}{p(\bar{x}|\theta_1)|B(\theta_1, \bar{x})|} - \Delta^P(\theta_0) \right). \end{aligned}$$

Therefore, a sufficient condition for $\lim_{\epsilon \rightarrow 0^+} \frac{dU^P[\Gamma^\epsilon]}{d\epsilon} > 0$ is that $\frac{\Delta^P(\theta)}{p(\bar{x}|\theta)|B(\theta, \bar{x})|}$ is strictly increasing over $[\theta_0, \theta_1]$. Property (3) in Condition (1) guarantees this is the case.

We conclude that, no matter whether, under the original policy $\tilde{\Gamma}$, $\hat{\theta}(\bar{x}) \geq \theta_1$, or $\hat{\theta}(\bar{x}) < \theta_1$, starting from $\tilde{\Gamma}$, there exist policies in $\mathbb{G}$ that strictly improve upon $\tilde{\Gamma}$, which contradicts the optimality of $\tilde{\Gamma}$. This means that any policy that maximizes the policy maker's payoff over $\mathbb{G}$ is such that there exists θ^* such that $\tilde{\pi}(0|\theta) = 0$ for F -almost all $\theta \leq \theta^*$ and $\tilde{\pi}(0|\theta) = 1$ for F -almost all $\theta > \theta^*$.

Now recall that the original policy $\Gamma = (\{0, 1\}, \pi)$ is such that there exists no θ^* such that $\pi(0|\theta) = 0$ for F -almost all $\theta \leq \theta^*$ and $\pi(0|\theta) = 1$ for F -almost all $\theta > \theta^*$. This means that any threshold policy $\tilde{\Gamma}$ that maximizes the policy maker's payoff over $\mathbb{G}$ yields a payoff strictly higher than Γ . The result in the theorem then follows from observing that, given $\tilde{\Gamma}$, there exists a nearby threshold policy $\Gamma^* \in \mathbb{G}$ with cut-off θ^{**} arbitrarily close θ^* that satisfies the perfect coordination

property (i.e., such that $U^{\Gamma^*}(x, 0|x) > 0$ all x) and that yields the policy maker a payoff arbitrarily close to that of $\tilde{\Gamma}$. To see this, it suffices to note that the threshold θ^* that defines $\tilde{\pi}$ in $\tilde{\Gamma}$ is given by

$$\theta^* \equiv \inf\{\theta' : \int_{\theta'}^{\infty} u(\theta, P(x|\theta))p(x|\theta)f(\theta)d\theta > 0 \text{ for all } x \in \mathbb{R}\}.$$

The policy maker's payoff under $\tilde{\Gamma}$ can then be approximated arbitrarily well by any threshold policy with cut-off equal to $\theta^* + \varepsilon$, for $\varepsilon > 0$ but small. Because any such policy satisfies the perfect coordination policy, we then have that the result in the theorem holds. Q.E.D.

Proof of Proposition 3. We establish the result by showing that, when Condition (7) holds, for any fixed $\hat{\theta}$, the function $\Psi(\hat{\theta}, \sigma_z) \equiv \min_{\theta_0} \psi(\theta_0, \hat{\theta}, \sigma_z)$ is increasing in σ_z . Moreover, in this case, the regime threshold in the absence of any public disclosure, $\theta_{\sigma_z}^*$, is decreasing in σ_z , and $\lim_{\sigma_z \rightarrow 0+} \theta_{\sigma_z}^* = \theta^{MS}$.

To ease the notation, let $\sigma = \sigma_z$. By the envelope theorem, we have that $\frac{\partial}{\partial \sigma} \Psi(\hat{\theta}, \sigma) = \frac{\partial}{\partial \sigma} \psi(\bar{\theta}_\sigma, \hat{\theta}, \sigma)$, with $\bar{\theta}_\sigma \in \arg \min_{\theta_0} \psi(\theta_0, \hat{\theta}, \sigma)$. Note that, for any $\theta_0 > \hat{\theta}$, any σ ,

$$\begin{aligned} \frac{\partial}{\partial \sigma} \psi(\theta_0, \hat{\theta}, \sigma) &= \frac{\partial}{\partial \sigma} \int_{\hat{\theta}}^{\infty} (\bar{b}(\theta) \mathbb{I}_{\theta < \theta_0} + \bar{g}(\theta) \mathbb{I}_{\theta \geq \theta_0}) \frac{\phi\left(\frac{z_{\sigma}^*(\theta_0) - \theta}{\sigma}\right)}{\sigma \Phi\left(\frac{z_{\sigma}^*(\theta_0) - \hat{\theta}}{\sigma}\right)} d\theta \\ &= \frac{\frac{\partial}{\partial \sigma} \int_{\hat{\theta}}^{\Phi\left(\frac{z_{\sigma}^*(\theta_0) - \hat{\theta}}{\sigma}\right)} (\bar{b}(z_{\sigma}^*(\theta_0) - \sigma \Phi^{-1}(A)) \mathbb{I}_{A > \theta_0} + \bar{g}(z_{\sigma}^*(\theta_0) - \sigma \Phi^{-1}(A)) \mathbb{I}_{A \leq \theta_0}) dA}{\Phi\left(\frac{z_{\sigma}^*(\theta_0) - \hat{\theta}}{\sigma}\right)} \\ &= \frac{\int_{\hat{\theta}}^{\Phi\left(\frac{z_{\sigma}^*(\theta_0) - \hat{\theta}}{\sigma}\right)} (\bar{b}'(z_{\sigma}^*(\theta_0) - \sigma \Phi^{-1}(A)) \mathbb{I}_{A > \theta_0} + \bar{g}'(z_{\sigma}^*(\theta_0) - \sigma \Phi^{-1}(A)) \mathbb{I}_{A \leq \theta_0}) (\Phi^{-1}(\theta_0) - \Phi^{-1}(A)) dA}{\Phi\left(\frac{z_{\sigma}^*(\theta_0) - \hat{\theta}}{\sigma}\right)} \\ &\quad + \frac{(\psi(\theta_0, \hat{\theta}, \sigma) - b(\hat{\theta})) \phi\left(\frac{z_{\sigma}^*(\theta_0) - \hat{\theta}}{\sigma}\right) (\theta_0 - \hat{\theta})}{\sigma^2 \Phi\left(\frac{z_{\sigma}^*(\theta_0) - \hat{\theta}}{\sigma}\right)} \end{aligned}$$

where the second equality follows from the change of variables $A = \Phi\left(\frac{z_{\sigma}^*(\theta_0) - \theta}{\sigma}\right)$ along with the fact that, by definition, $\Phi\left(\frac{z_{\sigma}^*(\theta_0) - \theta_0}{\sigma}\right) = \theta_0$, while the third equality from using $z_{\sigma}^*(\theta) = \theta + \sigma \Phi^{-1}(\theta)$. Lastly, by reverting the change of variables, and letting

$$D(\theta, \theta_0) \equiv \begin{cases} \bar{b}'(\theta) & \text{if } \theta < \theta_0 \\ \bar{g}'(\theta) & \text{if } \theta \geq \theta_0, \end{cases}$$

we have that

$$\begin{aligned} \frac{\partial}{\partial \sigma} \psi(\theta_0, \hat{\theta}, \sigma) &= \frac{\int_{\hat{\theta}}^{\infty} D(\theta, \theta_0) (\theta - \theta_0) \phi\left(\frac{z_{\sigma}^*(\theta_0) - \theta}{\sigma}\right) d\theta + (\psi(\theta_0, \hat{\theta}, \sigma) - b(\hat{\theta})) \phi\left(\frac{z_{\sigma}^*(\theta_0) - \hat{\theta}}{\sigma}\right) (\theta_0 - \hat{\theta})}{\sigma^2 \Phi\left(\frac{z_{\sigma}^*(\theta_0) - \hat{\theta}}{\sigma}\right)} \\ &= \sigma^{-1} \mathbb{E}[D(\theta, \theta_0) (\theta - \theta_0) | z_{\sigma}^*(\theta_0), \theta \geq \hat{\theta}] + \frac{(\psi(\theta_0, \hat{\theta}, \sigma) - b(\hat{\theta})) \phi\left(\frac{z_{\sigma}^*(\theta_0) - \hat{\theta}}{\sigma}\right) (\theta_0 - \hat{\theta})}{\sigma^2 \Phi\left(\frac{z_{\sigma}^*(\theta_0) - \hat{\theta}}{\sigma}\right)}. \end{aligned}$$

When evaluated at $\hat{\theta} = \theta_{\sigma}^{inf}$ and $\theta_0 = \theta_{\sigma}^{\#}$, because $\psi(\theta_{\sigma}^{\#}, \theta_{\sigma}^{inf}, \sigma) = 0$, we have that the above expression becomes

$$\frac{\partial}{\partial \sigma} \psi(\theta_{\sigma}^{\#}, \theta_{\sigma}^{inf}, \sigma) = \sigma^{-1} \mathbb{E}[D(\theta, \theta_{\sigma}^{\#})(\theta - \theta_{\sigma}^{\#}) | z_{\sigma}^*(\theta_{\sigma}^{\#}), \theta \geq \theta_{\sigma}^{inf}] + \frac{|b(\theta_{\sigma}^{inf})| \phi\left(\frac{z_{\sigma}^*(\theta_{\sigma}^{\#}) - \theta_{\sigma}^{inf}}{\sigma}\right) (\theta_{\sigma}^{\#} - \theta_{\sigma}^{inf})}{\sigma^2 \Phi\left(\frac{z_{\sigma}^*(\theta_{\sigma}^{\#}) - \theta_{\sigma}^{inf}}{\sigma}\right)}. \quad (21)$$

It is now easy to see that Condition (7) implies that $\frac{\partial}{\partial \sigma} \Psi(\theta_{\sigma}^{\#}, \theta_{\sigma}^{inf}, \sigma) > 0$.

The above property implies that, for given $\theta_{\sigma_z}^{inf}$, a marginal increase in the standard deviation of the agents' private information at σ_z increases $\Psi(\theta_{\sigma_z}^{inf}, \sigma_z)$. Furthermore, because the threshold $\theta_{\sigma_z}^{\#}$ solves $\psi(\theta_{\sigma_z}^{\#}, \theta_{\sigma_z}^{inf}, \sigma_z) = 0$, we have that, by increasing σ_z while keeping the threshold $\theta_{\sigma_z}^{inf}$ fixed, the policy maker guarantees that, for any $\theta > \theta_{\sigma_z}^{inf}$, $\psi(\theta, \theta_{\sigma_z}^{inf}, \sigma_z) > 0$. Next, note that $\theta_{\sigma_z}^{inf}$ is decreasing in σ_z . This follows from the fact that, for any σ_z , any $\theta > \hat{\theta}$, $\psi(\theta, \hat{\theta}, \sigma_z)$ is strictly increasing in $\hat{\theta}$ (this last property in turn follows from Lemma 2 in Angeletos et al. (2007)). From the above results, we thus have that, starting from any discriminatory policy, a reduction in the precision of the agents' private information (i.e., a marginal increase in σ_z) lowers the fundamental threshold $\theta_{\sigma_z}^{inf}$ below which regime change occurs, thus improving the policy maker's payoff. Q.E.D.

References

- Allen, F., Gale, D., 1998. Optimal financial crises. *Journal of Finance* 53, 1245–1284.
- Alonso, R., Camara, O., 2016a. Persuading voters. *American Economic Review* 106, 3590–3605.
- Alonso, R., Camara, O., 2016b. Bayesian persuasion with heterogeneous priors. *Journal of Economic Theory* 165, 672–706.
- Alvarez, F., Barlevy, G., 2015. Mandatory disclosure and financial contagion. WP, NBER .
- Angeletos, G.M., Hellwig, C., Pavan, A., 2006. Signaling in a global game: Coordination and policy traps. *Journal of Political Economy* 114, 452–484.
- Angeletos, G.M., Hellwig, C., Pavan, A., 2007. Dynamic global games of regime change: Learning, multiplicity, and the timing of attacks. *Econometrica* 75, 711–756.
- Angeletos, G.M., Pavan, A., 2013. Selection-free predictions in global games with endogenous information and multiple equilibria. *Theoretical Economics* 8, 883–938.
- Arieli, I., Babichenko, Y., 2017. Private bayesian persuasion. WP, Technion-Israel Institute of Technology .
- Bardhi, A., Guo, Y., 2017. Modes of persuasion toward unanimous consent. *Theoretical Economics*, forthcoming .
- Basak, D., Zhou, Z., 2017. Diffusing coordination risk. WP, NYU .
- Bergemann, D., Morris, S., 2016a. Bayes correlated equilibrium and the comparison of information structures in games. *Theoretical Economics* 11, 487–522.

- Bergemann, D., Morris, S., 2016b. Information design, bayesian persuasion and bayes correlated equilibrium. *The American Economic Review* 106, 586–591.
- Bergemann, D., Morris, S., 2017. Information design: A unified perspective. WP, Yale University .
- Bouvard, M., Chaigneau, P., Motta, A.d., 2015. Transparency in the financial system: Rollover risk and crises. *The Journal of Finance* 70, 1805–1837.
- Calzolari, G., Pavan, A., 2006a. On the optimality of privacy in sequential contracting. *Journal of Economic theory* 130, 168–204.
- Calzolari, G., Pavan, A., 2006b. Monopoly with resale. *The RAND Journal of Economics* 37, 362–375.
- Faria-e Castro, M., Martinez, J., Philippon, T., 2016. Runs versus lemons: information disclosure and fiscal capacity. *The Review of Economic Studies* , 060.
- Chan, J., Gupta, S., Li, F., Wang, Y., 2016. Pivotal persuasion. WP, NYU .
- Che, Y.K., Hörner, J., 2017. Recommender systems as mechanisms for social learning. *Quarterly Journal of Economics* .
- Cong, L.W., Grenadier, S., Hu, Y., et al., 2016. Dynamic Coordination and Intervention Policy. Technical Report.
- Corona, C., Nan, L., Gaoqing, Z., 2017. The coordination role of stress test disclosure in bank risk taking. WP, Carnegie Mellon University .
- Denti, T., 2015. Unrestricted information acquisition. WP, Cornell University .
- Doval, L., Ely, J., 2017. Sequential information design. WP, Northwestern University .
- Dworczak, P., 2016. Mechanism design with aftermarkets: Cutoff mechanisms. WP, Stanford .
- Edmond, C., 2013. Information manipulation, coordination, and regime change. *Review of Economic Studies* 80, 1422–1458.
- Ely, J.C., 2017. Beeps. *The American Economic Review* 107, 31–53.
- Farhi, E., Tirole, J., 2012. Collective moral hazard, maturity mismatch, and systemic bailouts. *American Economic Review* 102, 60–93.
- Flannery, M., Hirtleb, B., Kovner, A., 2017. Evaluating the information in the federal reserve stress tests. *Journal of Financial Intermediation* 29, 1–18.
- Garcia, F., Panetti, E., 2017. A theory of government bailouts in a heterogeneous banking union. WP, Indiana University .
- Goldstein, I., Huang, C., 2016. Bayesian persuasion in coordination games. *The American Economic Review* 106, 592–596.
- Goldstein, I., Leitner, Y., 2015. Stress tests and information disclosure. WP, FRB of Philadelphia .
- Goldstein, I., Sapra, H., 2014. Should banks’ stress test results be disclosed? an analysis of the costs and benefits. *Foundations and Trends (R) in Finance* 8, 1–54.
- Henry, J., Christoffer, K., 2013. A macro stress testing framework for assessing systemic risks in the banking sector. WP, European Central Bank .

- Homar, T., Kick, H., Salleo, C., 2016. Making sense of the **EU** wide stress test: a comparison with the srisk approach. WP, European Central Bank .
- Iachan, F.S., Nenov, P.T., 2015. Information quality and crises in regime-change games. *Journal of Economic Theory* 158, 739–768.
- Kamenica, E., Gentzkow, M., 2011. Bayesian persuasion. *American Economic Review* 101, 2590–2615.
- Kolotilin, A., Mylovanov, T., Zapechelnyuk, A., Li, M., 2017. Persuasion of a privately informed receiver. *Econometrica* 85, 1949–1964.
- Leitner, Y., Williams, B., 2017. Model secrecy and stress tests. WP, Philadelphia Fed .
- Lerner, J., Tirole, J., 2006. A model of forum shopping. *American economic review* 96, 1091–1113.
- Mathevet, L., Perego, J., Taneva, I., 2016. Information design: The epistemic approach. WP, NYU .
- Mensch, J., 2017. Monotone persuasion. WP, Hebrew University .
- Morgan, D., Persitani, S., Vanessa, S., 2014. The information value of the stress test. *Journal of Money, Credit and Banking*, Vol. 46, No. 7 .
- Moriya, F., Yamashita, T., 2017. Asymmetric information allocation to avoid coordination failure. WP, Toulouse School of Economics .
- Morris, S., Shin, H.S., 2006. Global games: Theory and applications. *Advances in Economics and Econometrics* , 56.
- Morris, S., Yang, M., 2016. Coordination under continuous choice. WP, Princeton .
- Myerson, R.B., 1986. Multistage games with communication. *Econometrica* , 323–358.
- Orlov, D., Zryumov, P., Skrzypacz, A., 2017. Design of macro-prudential stress tests. WP, Stanford University .
- Petrella, G., Resti, A., 2013. Supervisors as information producers: Do stress tests reduce bank opaqueness? *Journal of Banking and Finance* 37, 5406–5420.
- Segal, I., 2003. Coordination and discrimination in contracting with externalities: Divide and conquer? *Journal of Economic Theory* 113, 147–181.
- Shimoji, M., 2017. Bayesian persuasion in unlinked games. WP, University of York .
- Szkup, M., Trevino, I., 2015. Information acquisition in global games of regime change. *Journal of Economic Theory* 160, 387–428.
- Taneva, I.A., 2016. Information design. WP, University of Edinburgh .
- Wang, Y., 2015. Bayesian persuasion with multiple receivers. WP, University of Pittsburgh .
- Williams, B., 2017. Stress tests and bank portfolio choice. WP, New York University .
- Yang, M., 2015. Coordination with flexible information acquisition. *Journal of Economic Theory* 158, 721–738.