

Inostroza, Nicolás; Pavan, Alessandro

Article

Adversarial coordination and public information design

Theoretical Economics

Provided in Cooperation with:

The Econometric Society

Suggested Citation: Inostroza, Nicolás; Pavan, Alessandro (2025) : Adversarial coordination and public information design, Theoretical Economics, ISSN 1555-7561, The Econometric Society, New Haven, CT, Vol. 20, Iss. 2, pp. 763-813, <https://doi.org/10.3982/TE5768>

This Version is available at:

<https://hdl.handle.net/10419/320298>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.



<https://creativecommons.org/licenses/by-nc/4.0/>

Adversarial coordination and public information design

NICOLAS INOSTROZA

Rotman School of Management, University of Toronto and FTG

ALESSANDRO PAVAN

Department of Economics, Northwestern University and CEPR

We study flexible public information design in global games. In addition to receiving public information from the designer, agents are endowed with exogenous private information and must decide between two actions (invest and not invest), the profitability of which depends on unknown fundamentals and the agents' aggregate action. The designer does not trust the agents to play favorably to her and evaluates any policy under the "worst-case scenario." First, we show that the optimal policy removes any strategic uncertainty by inducing all agents to take the same action, but without permitting them to perfectly learn the fundamentals and/or the beliefs that rationalize other agents' actions. Second, we identify conditions under which the optimal policy is a simple "pass/fail" test. Finally, we show that when the designer cares only about the probability the aggregate investment is successful, the optimal policy need not be monotone in fundamentals but then identify conditions on payoffs and exogenous beliefs under which the optimal policy is monotone.

KEYWORDS. Global games, adversarial coordination, Bayesian persuasion, robust public information design.

JEL CLASSIFICATION. D83, G28, G33.

1. INTRODUCTION

Coordination plays a major role in many socioeconomic environments. The damages to society of miscoordination can be severe and often call for government intervention. Think of the possibility of default by major financial institutions in case investors run or refrain from rolling over their short-term positions. Such defaults can trigger a collapse in financial markets, with severe consequences for the real economy. Confronted with such prospects, governments and supervising authorities have incentives to intervene.

Nicolas Inostroza: Nicolas.inostroza@rotman.utoronto.ca

Alessandro Pavan: alepavan@northwestern.edu

A previous version circulated under the title "Persuasion in Global Games with Application to Stress Testing." For comments and useful suggestions, we thank the referees, Marios Angeletos, Gadi Barlevy, Eddie Dekel, Tommaso Denti, Laura Doval, Stephen Morris, Marciano Siniscalchi, Bruno Strulovici, Jean Tirole, Xavier Vives, Leifu Zhang, and seminar participants at various conferences and institutions where the paper was presented. Matteo Camboni provided excellent research assistance. Pavan acknowledges financial support from the NSF under the grant SES-1730483 and from Bocconi University where part of this research was conducted. The usual disclaimer applies.

These interventions often take the form of public information disclosures, such as stress testing or, more broadly, releases of information aimed at influencing market beliefs.

In this paper, we study *public information design* in markets in which a large number of receivers (e.g., investors in financial markets) must choose whether to play an action favorable to the designer (e.g., pledging funds to a financial institution), or an “adversarial” action (e.g., refraining from pledging). A policymaker can flexibly design a policy disclosing information to market participants about relevant economic fundamentals. The analysis delivers results that are important for various situations in which coordination plays a major role, including bank runs, currency crises, and technology and standards adoption. In the context of stress testing, the policymaker may represent a supervising authority attempting to prevent a run against the banking sector (see, e.g., [Henry and Kok \(2013\)](#) and [Homar, Kick, and Salleo \(2016\)](#)). In the case of currency crises, the policymaker may represent a central bank attempting to dissuade speculators from short-selling the domestic currency by releasing information about the bank’s reserves and/or domestic economic fundamentals. In the case of technology adoption, the policymaker may represent the owners of an intellectual property trying to persuade heterogeneous market users of the merits of a new product ([Lerner and Tirole \(2006\)](#)).

The backbone of the model is a *global game of regime change* in which multiple agents must choose between “attacking” a status quo or “refraining from attacking,” and where the success of the attack depends on its aggregate size and on exogenous fundamentals. In addition to receiving public information from the designer, agents are endowed with exogenous private information. The designer does not trust the agents to play favorably to her and evaluates any policy of her choice under the “worst-case” scenario. That is, when multiple rationalizable strategy profiles are consistent with the information disclosed, the designer takes a “robust approach” by looking at the outcome that prevails when agents play according to the rationalizable profile least favorable to her.¹

We assume the policymaker can flexibly design a policy that disseminates publicly information about relevant economic fundamentals. We use the model to address the following questions: (a) Are there benefits to preventing market participants from predicting each others’ actions and beliefs? (b) When are simple policies such as pass/fail tests optimal? (c) Are there merits to nonmonotone rules that induce the market to play favorably for intermediate fundamentals but not necessarily for stronger ones?

Our first result establishes that, despite the fear of adversarial coordination, the optimal policy satisfies the “*perfect coordination property*.” In each state, it induces all market participants to take the same action, but without creating homogenous beliefs among market participants. In other words, the optimal policy completely removes any *strategic uncertainty* while preserving *structural uncertainty*. Given the public information disclosed, each receiver can perfectly predict the action of any other receiver, but not the beliefs that rationalize such actions. For example, an agent who is induced to invest must not be able to determine whether other agents invest because they know that

¹Such a robust approach is motivated by the applications the analysis is meant for. For example, when concerned about runs to the banking sector, policymakers typically do not trust the market to play favorably.

the fundamentals are so strong that the investment will always succeed, irrespective of the behavior of other agents (e.g., the bank will never default), or because they are confident that other agents will invest. The optimal policy leverages the heterogeneity of the agents' primitive beliefs by making investing dominant for some agents based on their first-order beliefs, but only iteratively dominant for others based on their higher-order beliefs.² Under adversarial coordination, preserving uncertainty over beliefs is key to the minimization of the risk of an undesirable outcome such as a bank default, a currency collapse, or the failure of a new technology to take off. When the designer trusts the agents to follow her recommendations, the optimality of the perfect coordination property is straightforward and follows from arguments similar to those establishing the Revelation Principle (e.g., Myerson (1986)). This is not the case under adversarial coordination for information that facilitates perfect coordination may also favor the emergence of rationalizable profiles in which some of the agents play adversarially to the designer.

Our second result shows that, when the economic fundamentals and the agents' beliefs comove in the sense that states in which fundamentals are strong are also states in which most agents expect other agents to expect the fundamentals to be strong, and so on, then the optimal policy takes the form of a simple "pass/fail" test, with no further information disclosed to the market. It is known that, when the distribution from which the agents' private signals are drawn is log-supermodular, or equivalently, satisfies the *Monotone Likelihood Ratio Property*—in short MLRP—all agents follow monotone (i.e., cut-off) strategies, no matter the public information. This is because, under MLRP, the agents' "optimism ranking" is preserved under Bayesian updating. If agent j is more optimistic than agent i before the public announcement is made (formally, j 's beliefs dominate i 's beliefs according to the MLRP order), then this continues to be the case after *any* public announcement. When this is the case, disclosing information to the market in addition to whether or not the policymaker expects the agents' investment to succeed when they play adversarially does not help. We also show that MLRP is key to the optimality of simple pass/fail policies. When the information the policymaker discloses can be used to change the ranking of the agents' optimism, the policymaker can leverage the optimism reversal to spare more fundamentals from the undesirable outcome by disclosing information in addition to whether or not she expects the investment to succeed.³

In the context of stress testing, these results provide a foundation for the optimality of simple pass/fail policies. Importantly, optimal stress tests should be *transparent*, in the sense of facilitating coordination among investors, but should not generate consensus among market participants about the soundness of the financial institutions under scrutiny.

²The optimal policy does not ensure that investing is the unique rationalizable action based on first-order beliefs for all agents. It relies on a contagion argument *through higher-order beliefs* to induce all agents to invest under the unique rationalizable profile.

³When, instead, the designer trusts her ability to coordinate the receivers on the course of action most favorable to her, optimal policies always take the form of action recommendations, and hence pass/fail policies are optimal, irrespective of the agents' primitive beliefs. This is not the case under adversarial/robust design.

Our third result is about the optimality of monotone pass/fail policies, that is, rules that grant a pass grade if and only if the exogenous fundamentals are above a given threshold. We show that the optimality of such rules is related to the extent to which the policymaker's preferences for a favorable outcome (e.g., for avoiding a bank default) vary with the fundamentals. We identify precise conditions involving the policymaker's preferences and the agents' payoffs and exogenous beliefs under which monotone rules are optimal. Such conditions are fairly sharp in the sense that, when violated, one can identify instances in which nonmonotone rules strictly outperform monotone ones.⁴ The reason is that nonmonotone rules make it more difficult for the agents to *commonly learn* the fundamentals and hence permit the policymaker to give a pass grade to a larger set of fundamentals. When the policymaker's preferences for the favorable outcome (i.e., for avoiding a bank default) do not vary much with the exogenous fundamentals (in particular, when they are constant), nonmonotone rules may be optimal.

Organization 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 contains all the results about properties of optimal policies (perfect-coordination, pass/fail, monotonicity). Section 4 discusses how the results accommodate enrichments that are useful in applications (e.g., more general payoffs, as well as the possibility that the policymaker faces uncertainty about the outcome induced by her information dissemination). Section 5 concludes. The Appendix contains all proofs with the exception of the proofs of Examples 2 and 3, which are in the Supplementary Appendix (available at <https://econtheory.org/supp/5768/supplement.pdf>). The manuscript Inostroza and Pavan (2024a) contains additional material. In particular, (a) it extends Theorem 1* in the main text (about the optimality of perfectly coordinating the market response) to a broader class of economies, (b) discusses the benefits of discriminatory disclosures, when the latter are feasible, and (c) expands the material in Section 4.3 discussing the role of the multiplicity of the receivers and their exogenous private information for the optimality of monotone rules.

(Most) pertinent literature The paper is related to a few strands of the literature. The first one is the literature on *adversarial coordination and unique implementation*. See, among others, Segal (2003), Winter (2004), Sakovics and Steiner (2012), Frankel (2017), Halac, Kremer, and Winter (2020), and Halac, Lipnowski, and Rappoport (2021). These papers focus on the design of transfers. Instead, we focus on the design of public information in settings in which the receivers are endowed with exogenous private information. Li, Song, and Zhao (2023), and Morris, Oyama, and Takahashi (2024) consider the design of *private* information in binary supermodular games in which the receivers' exogenous information is symmetric. Halac, Lipnowski, and Rappoport (2022) study unique implementation when the designer can use a combination of transfers and information provision. Goldstein and Huang (2016) and Galvão and Shalders (2022) also

⁴We also show that the conditions guaranteeing the optimality of monotone rules are more stringent when the policymaker faces multiple privately-informed receivers than when she faces either a single (possibly privately-informed) receiver, or multiple receivers who possess no exogenous private information.

study public information design in settings in which the receivers possess exogenous private information. Goldstein and Huang (2016) restrict the policymaker to binary monotone rules, whereas Galvão and Shalders (2022) to partitional structures whereby when two states are pooled into the same cell, all in-between states are also pooled into the same cell. Related is also Alonso and Zachariadis (2023) who study the complementarity between private and public information. Relative to these works, our paper establishes three key results: (a) it proves that inducing all agents to take the same action is always optimal, despite the fear of adversarial coordination; (b) it shows why, in general, binary policies are suboptimal but then identifies sharp conditions under which such policies are optimal; (c) it shows why, in general, nonmonotone rules permit the policymaker to induce a favorable outcome over a larger set of fundamentals but then identifies sharp conditions under which optimal policies are monotone.⁵

The second strand is the literature on *information design with multiple receivers*. See, among others, Alonso and Camara (2016a), Arieli and Babichenko (2019), Bardhi and Guo (2018), Basak and Zhou (2020), Che and Hörner (2018), Doval and Ely (2020), Galperti and Perego (2023), Gick and Pausch (2012), Gitmez and Molavi (2022), Heese and Lauermann (2021), Laclau and Renou (2017), Mathevet, Perego, and Taneva (2020), Shimoji (2021), and Taneva (2019). The key contribution vis-a-vis this literature is in showing how the interaction between (a) adversarial coordination and (b) exogenous private information among the receivers shapes the optimal provision of public information.⁶

The third strand is the literature on *global games with endogenous information*. Angeletos, Hellwig, and Pavan (2006) and Angeletos and Pavan (2013) study signaling in global games. Angeletos and Werning (2006) investigate the role of prices as a vehicle for information aggregation. Angeletos, Hellwig, and Pavan (2007) consider a dynamic model in which agents learn from the accumulation of private information and from the (possibly noisy) observation of past outcomes. Cong, Grenadier, and Hu (2020) consider a dynamic setting similar to the one in Angeletos, Hellwig, and Pavan (2007) but allowing for policy interventions. Edmond (2013) and Kyriazis and Lou (2024) consider propaganda in global games, in a setting in which the policymaker manipulates the agents' private signals. Szkup and Trevino (2015), Yang (2015), Morris and Yang (2022), and Denti (2023) study the acquisition of private information in global games. Our paper contributes to this strand by identifying properties of flexible public information provision when (a) the sender can commit, and (b) the receivers play adversarially.

Finally, the paper is related to the literature on *stress testing*. See Goldstein and Sapra (2014) for an overview of some of the early contributions. Bouvard, Chaigneau, and de Motta (2015) study a setting where a policymaker must choose between full transparency and full opacity but cannot commit to a disclosure policy. Williams (2017) and Goldstein and Leitner (2018) study the design of stress tests when the receivers do not possess exogenous private information. Orlov, Zryumov, and Skrzypacz (2023) study the joint design of stress tests and precautionary recapitalizations whereas Faria-e Castro,

⁵In particular, Example 3 below shows that nonmonotone rules strictly outperform monotone ones in the same environment of Goldstein and Huang (2016).

⁶See Bergemann and Morris (2019) and Kamenica (2019) for an overview on information design.

Martinez, and Philippon (2016) and Garcia and Panetti (2017) the joint design of stress tests and government bailouts. Inostroza (2023) studies regulatory disclosures with multiple audiences of investors who care about different aspects of a financial institution's balance sheet. Alvarez and Barlevy (2021) and Quigley and Walther (2024) study the incentives of banks to disclose balance sheet (hard) information. Corona, Nan, and Gao-qing (2017) study how stress tests disclosures may favor banks' coordinated risk taking in the spirit of Farhi and Tirole (2012). Morgan, Persitani, and Vanessa (2014), Flannery, Hirtle, and Kovner (2017), and Petrella and Resti (2013) conduct an empirical analysis of the information provided by stress tests in the US and the EU. Our paper contributes to this literature along the following dimensions: (a) it shows that optimal stress tests should not create conformism in market participants' beliefs about exogenous fundamentals but should be sufficiently transparent to eliminate any ambiguity about the market response to the tests; (b) it identifies conditions under which simple pass/fail policies are optimal; (c) it provides conditions for optimal tests to be monotone (see also Inostroza and Pavan (2024b) for a discussion of how the toughness of optimal stress tests relates to the type of securities issued by the banks).

2. MODEL

Global games have been used to study the interaction between information and coordination in many socioeconomic environments, including bank runs, debt crises, currency attacks, investment in technologies with network externalities, technological spillovers, and political change.

To ease the exposition, hereafter we describe the model and all the results in the context of a specific game in the spirit of Rochet and Vives (2004) in which the agents are investors (e.g., fund managers, or unsecured bank depositors) deciding whether or not to pledge funds to one, or multiple financial institutions, and where these institutions default on their obligations when the size of the aggregate investment is not large enough.⁷ The analysis, however, readily extends to many other global games.

Players and actions A policymaker designs an information disclosure policy, for example, stress tests, call reports, publication of accounting standards, and disclosure of various macro and financial variables that are jointly responsible for the profitability of the agents' decisions. The market is populated by a measure-one continuum of agents (the receivers) distributed uniformly over $[0, 1]$. Each agent may either take a "friendly" action, $a_i = 1$, or an "adversarial" action, $a_i = 0$. The friendly action is interpreted as the decision to invest (more generally, to "refrain from attacking" a status quo the policymaker wants to preserve). The adversarial action is interpreted as the decision to not invest (more generally, to "attack"). We denote by $A \equiv \int_0^1 a_i di \in [0, 1]$ the size of the aggregate investment.

⁷Rochet and Vives (2004) consider a three-period economy à la Diamond and Dybvig (1983) but with heterogeneous investors, in which banks may fail early or late. As shown in that paper, the full model admits a reduced-form version similar to the one considered here.

Fundamentals and exogenous information Consistently with the rest of the literature, we parameterize the relevant fundamentals by $\theta \in \mathbb{R}$. The fundamentals are exogenous to the policymaker's choice of a disclosure policy. It is commonly believed (by the policymaker and the agents alike) that θ is drawn from a distribution F , absolutely continuous over an interval $\Theta \supsetneq [0, 1]$, with a smooth density f strictly positive over Θ . In addition, each agent $i \in [0, 1]$ is endowed with private information summarized by a unidimensional statistic $x_i \in \mathbb{R}$ drawn independently across agents given θ from an absolutely continuous cumulative distribution function $P(x|\theta)$ with smooth density $p(x|\theta)$ strictly positive over an (open) interval $\varrho_\theta \equiv (\underline{\varrho}_\theta, \bar{\varrho}_\theta)$ containing θ , with $\underline{\varrho}_\theta, \bar{\varrho}_\theta$ monotone in θ , and with $p(x|\theta)$ bounded over (x, θ) . The bounds $\underline{\varrho}_\theta, \bar{\varrho}_\theta$ can be either finite or infinite. For example, when $x_i = \theta + \sigma \varepsilon_i$, with ε_i drawn from a uniform distribution over $[-1, +1]$, then for any θ , $\underline{\varrho}_\theta = \theta - \sigma$ and $\bar{\varrho}_\theta = \theta + \sigma$. When, instead, $x_i = \theta + \sigma \varepsilon_i$, with ε_i drawn from a standard Normal distribution, then for any θ , $\underline{\varrho}_\theta = -\infty$ and $\bar{\varrho}_\theta = +\infty$. Furthermore, in this latter case, $P(x|\theta) = \Phi((x - \theta)/\sigma)$, where Φ is the cumulative distribution function of the standard Normal distribution. We denote by $\mathbf{x} \equiv (x_i)_{i \in [0, 1]}$ a profile of private signals and by $\mathbf{X}(\theta)$ the collection of all $\mathbf{x} \in \mathbb{R}^{[0, 1]}$ that are consistent with the fundamentals being equal to θ . As usual, we assume that any pair of signal profiles $\mathbf{x}, \mathbf{x}' \in \mathbf{X}(\theta)$ has the same cross-sectional distribution of signals, with the latter equal to $P(x|\theta)$.

Regime outcome The fundamentals θ parameterize the critical size of the aggregate investment that is necessary to avoid default (more generally, an undesirable regime change). If $A > 1 - \theta$, short-term obligations are met and default is avoided. If, instead, $A \leq 1 - \theta$, default occurs. We denote by $r = 1$ the event in which default is avoided and by $r = 0$ the event in which default occurs.⁸

Dominance regions For any $\theta \leq 0$, default occurs irrespective of the size of the aggregate investment, whereas for any $\theta > 1$ default is averted with certainty. For $\theta \in (0, 1]$, instead, whether or not default occurs is determined by the behavior of the market.

Payoffs Each agent's payoff differential between investing and not investing, $u(\theta, A)$, is equal to $g(\theta) > 0$ in case default is avoided, and $b(\theta) < 0$ otherwise. In turn, the policymaker's payoff is equal to $W(\theta)$ in case default is avoided, and $L(\theta)$ in case of default, with $W(\theta) > L(\theta)$ for all θ . When W and L are invariant in θ , the policymaker's objective reduces to minimizing the probability of default. The functions b, g, W , and L are all bounded. For any $(\theta, A) \in \Theta \times [0, 1]$, then let

$$\begin{aligned} u(\theta, A) &\equiv g(\theta)\mathbf{1}(A > 1 - \theta) + b(\theta)\mathbf{1}(A \leq 1 - \theta), \\ U^P(\theta, A) &\equiv W(\theta)\mathbf{1}(A > 1 - \theta) + L(\theta)\mathbf{1}(A \leq 1 - \theta) \end{aligned}$$

denote the payoffs of a representative agent and of the policymaker, respectively, when the fundamentals are θ and the aggregate investment is A .

⁸The model assumes that, given A and θ , the regime outcome is binary. The case in which default is "partial" is qualitatively similar, from a strategic standpoint, to the case where, given A and θ , the regime outcome is stochastic and determined by variables that are not observable by the policymaker at the time of her public announcements (see the discussion in Section 4).

Policy Let $\mathcal{S}$ be a compact Polish space defining the set of possible signal realizations. A *policy* $\Gamma = (\mathcal{S}, \pi)$ consists of the set $\mathcal{S}$ along with a measurable mapping $\pi : \Theta \rightarrow \Delta(\mathcal{S})$ specifying, for each θ , a probability distribution over the information disclosed to the market.

Timing The sequence of events is the following:

- (i) The policymaker publicly announces the policy $\Gamma = (\mathcal{S}, \pi)$ and commits to it.⁹
- (ii) The fundamentals θ are drawn from the distribution F and the agents' exogenous signals $\mathbf{x} \in \mathbf{X}(\theta)$ are drawn from the distribution $P(x|\theta)$.
- (iii) The public signal s is drawn from the distribution $\pi(\theta)$ and is publicly observed.
- (iv) Agents simultaneously choose whether or not to invest.
- (v) The regime outcome is determined (i.e., whether or not default occurred) and payoffs are realized.

Adversarial coordination and robust information design The policymaker does not trust the market to follow her recommendations and play favorably to her (i.e., invest whenever $\theta > 0$).¹⁰ Instead, she adopts a robust/conservative approach. She evaluates any policy Γ under the “worst-case” scenario, that is, she assumes that the market plays according to the rationalizable strategy profile that is most adversarial to her, among all those consistent with the policy Γ .

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 policymaker's ex ante expected payoff over all profiles surviving *iterated deletion of interim strictly dominated strategies* (henceforth IDISDS).

In the IDISDS procedure leading to MARP, agents use Bayes rule to update their beliefs about the fundamentals θ and the other agents' exogenous information $\mathbf{x} \in \mathbf{X}(\theta)$ using the common prior F , the distribution of private signals $P(x|\theta)$, and the policy Γ . Under MARP, given (x, s) , each agent $i \in [0, 1]$, after receiving exogenous information x from Nature and endogenous information s from the policymaker, refrains from investing whenever there exists at least one conjecture over (θ, A) consistent with the above Bayesian updating and supported by all other agents playing strategies surviving IDISDS, under which refraining from investing is a best response for the individual.

REMARKS. Hereafter, we confine attention to policies Γ for which MARP exists.¹¹ Because the game among the agents is supermodular (no matter the prior F , the distribution P from which the exogenous signals are drawn, and the policy Γ), the strategy

⁹See Leitner and Williams (2023) for a discussion of the commitment assumption in stress testing.

¹⁰If she did, a simple monotone policy revealing whether or not $\theta > 0$ would be optimal.

¹¹Because the state is continuous, in principle, one can think of policies Γ for which the agents' common posteriors are not well-defined or, when combined with the agents' exogenous information, are such that the agents' hierarchies of beliefs are not well-defined, in which case MARP may not exist.

profile a^Γ coincides with the “smallest” Bayes–Nash equilibrium (BNE) of the continuation game among the agents, and minimizes the policymaker’s payoff state by state, and not just in expectation. The reason why we consider MARP is that, in general, without imposing specific assumptions on F , P , and Γ , the only way the “smallest” BNE can be identified is by the iterated deletion of interim strictly dominated strategies. In standard global games, the “smallest” BNE is typically identified by assuming the agents’ signals are drawn from a distribution P satisfying the monotone likelihood property (MLRP), which is also used to guarantee equilibrium uniqueness. Here, we allow for arbitrary policies Γ , and do not require that, given Γ , the continuation equilibrium be unique.

Furthermore, given a policy $\Gamma = (\mathcal{S}, \pi)$, when describing the agents’ behavior, we do not distinguish between pairs (x, s) that are mutually consistent given Γ (meaning that the joint density of (x, s) is positive, that is, $\int_{\theta: s \in \text{supp}(\pi(\theta))} p(x|\theta) dF(\theta) > 0$) and those that are not. Because the policymaker commits to the policy Γ , the abuse is legitimate and permits us to ease the exposition. Any claim about the optimality of the agents’ behavior, however, should be interpreted to apply to pairs (x, s) that are mutually consistent given Γ .

3. PROPERTIES OF OPTIMAL POLICIES

We now introduce and discuss three key properties of optimal policies.

3.1 Perfect-coordination property

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

A disclosure policy thus has the perfect-coordination property if it coordinates all market participants on the same action, after any information it discloses. For any $\theta \in \Theta$, any $s \in \text{supp}(\pi(\theta))$, let $r^\Gamma(\theta, s) \in \{0, 1\}$ denote the regime outcome that prevails when agents play according to a^Γ , that is, $r^\Gamma(\theta, s) = 1$ (alternatively, $r^\Gamma(\theta, s) = 0$) means that default does not occur (alternatively, occurs) when, given (θ, s) , market participants play according to MARP consistent with Γ . That the agents’ signals are drawn independently from $P(x|\theta)$, conditional on θ , implies that the cross-sectional distribution of signals is pinned down by $P(x|\theta)$, and hence the regime outcome (i.e., whether default occurs or not) is the same across any pair of signal profiles $\mathbf{x}, \mathbf{x}' \in \mathbf{X}(\theta)$, and thus depends only on Γ , θ , and s . Hereafter, we say that the policy Γ is *regular* if MARP under Γ is well-defined and the regime outcome under a^Γ is measurable in (θ, s) .

THEOREM 1. *Given any regular policy Γ , there exists another regular policy Γ^* satisfying the perfect-coordination property (PCP) and such that, when the agents play according to MARP under both Γ and Γ^* , for any θ , (a) the probability of default under Γ^* is the same as under Γ , (b) the transition from Γ to Γ^* leads to a Pareto improvement (the policymaker is indifferent, no agent is worse off, and some agents are strictly better off).*

The policy Γ^* is obtained from the original policy Γ by disclosing, for each θ , in addition to the information $s \in \text{supp}(\pi(\theta))$ disclosed by the original policy Γ , a second piece of information that reveals to the market the regime outcome $r^\Gamma(\theta, s) \in \{0, 1\}$ that prevails at (θ, s) when agents play according to MARP consistent with the original policy Γ , a^Γ .

That, under the new policy Γ^* , it is rationalizable for all agents to invest when the policy discloses the information $(s, r^\Gamma(\theta, s)) = (s, 1)$, and to refrain from investing when the policy discloses the information $(s, r^\Gamma(\theta, s)) = (s, 0)$, is fairly straightforward. In fact, the announcement of $(s, 1)$ (alternatively, of $(s, 0)$) makes it common certainty among the agents that $\theta > 0$ (alternatively, that $\theta \leq 1$).

The reason why the result is not obvious is that the designer does not content herself with one rationalizable profile delivering the desired outcome; she is concerned with the possibility of adversarial coordination and, as a result, when she recommends to all the agents to invest, she must guarantee that investing is the *unique* rationalizable action for each agent, irrespective of his exogenous signal x . The proof in the [Appendix](#) shows that when the additional information is $r^\Gamma(\theta, s)$, this is indeed the case.

To fix ideas, consider first the case where, under the original policy Γ , the regime outcome $r^\Gamma(\theta, s)$ is monotone in θ . The announcement that $r^\Gamma(\theta, s) = 1$ makes it common certainty among the agents that $\theta > \hat{\theta}(s)$, for some threshold $\hat{\theta}(s)$. In this case, all agents revise their first-order beliefs about θ upward when receiving the additional information that $r^\Gamma(\theta, s) = 1$. That each agent is more optimistic about the strength of the fundamentals, however, does not guarantee that, under MARP consistent with the new policy Γ^* , more agents invest than under the original policy Γ . In fact, the new piece of information changes not only the agents' first-order beliefs about θ but also their higher-order beliefs and the latter matter for the determination of the most-aggressive rationalizable profile. More generally, $r^\Gamma(\theta, s)$ need not be monotone in θ . This is because MARP under the original policy Γ need not entail strategies that are monotone in x . As a result, in general, the announcement that $r^\Gamma(\theta, s) = 1$ need not trigger an upward revision of the agents' beliefs.

The result in Theorem 1 follows instead from the game being supermodular along with the fact that Bayesian updating preserves the likelihood ratio of any two states that are consistent with no default under the original policy Γ . Formally, for any $s \in \text{supp}(\pi(\Theta))$, any pair of states θ' and θ'' such that (a) $s \in \text{supp } \pi(\theta') \cap \text{supp } \pi(\theta'')$, and (b) $r^\Gamma(\theta', s) = r^\Gamma(\theta'', s) = 1$, the likelihood ratio of such two states under Γ^* is the same as under the original policy Γ . This property implies that the posterior beliefs (over Θ) of each agent with private signal x who, under the new policy Γ^* , receives information $(s, 1)$, are a “truncation” of the posterior beliefs the same agent would have had under the original policy Γ after receiving information s . The truncation eliminates from the support of the agent's original beliefs states θ at which, under MARP consistent with the original policy Γ there would have been default, and hence the agent's payoff differential from investing would have been negative. Because the game is supermodular, under any policy Γ , MARP is less aggressive than the most aggressive strategy profile surviving $n - 1$ rounds of IDISDS (in the sense that any agent who invests under the latter profile does so also under MARP, but the opposite is not necessarily true). This means that

the extra information $r^\Gamma(\theta, s) = 1$ also removes from the support of each agent's beliefs states θ at which the payoff differential from investing is negative under the most aggressive profile surviving $n - 1$ rounds of IDISDS under Γ . Hence, at any stage n of the IDISDS procedure, the truncation makes each agent more willing to invest. That is, any agent who would have invested after hearing s under the original policy Γ , also invests after hearing $(s, 1)$ under the new policy Γ^* . Because this is true for any n , it is also true in the limit as n goes to infinity. In other words, after the new policy Γ^* announces $(s, 1)$, each agent learns that his payoff differential from investing when all other agents play according to MARP consistent with the new policy Γ^* is strictly positive. Hence, after the new policy announces $(s, 1)$, each agent's unique rationalizable action is to invest, irrespective of her private information x .

When, instead, the new policy Γ^* announces $(s, 0)$, each agent learns that the state θ is among those at which there would have been default under MARP consistent with the original policy Γ (i.e., $r^\Gamma(\theta, s) = 0$). The announcement thus makes it common certainty among the agents that $\theta \leq 1$. It is then immediate that, under MARP consistent with the new policy Γ^* , all agents refrain from investing.

The policy Γ^* thus completely removes any strategic uncertainty. Indeed, when $(s, r^\Gamma(\theta, s)) = (s, 1)$ (alternatively, $(s, r^\Gamma(\theta, s)) = (s, 0)$) is announced, each agent knows that, under MARP consistent with the new policy Γ^* , all other agents invest (alternatively, refrain from investing), irrespective of their exogenous private information. Importantly, while the policy Γ^* removes any strategic uncertainty, it preserves structural uncertainty, that is, heterogeneity in the agents' first and higher-order beliefs about θ . As explained in the [Introduction](#), it is essential that agents who invest are uncertain as to whether other agents invest because they find it dominant to do so, or because when they count on other agents investing, they find it iteratively dominant to do so, which requires heterogeneity in posterior beliefs.

That the policymaker is indifferent between Γ and Γ^* is a direct implication of the fact that, for any θ , her payoff depends on $\mathcal{A}$ only through the probability of default, which is the same across the two policies. That, for any θ , no agent is worse off (and some agents are strictly better off) follows from the fact that, under Γ^* , all agents refrain from investing (alternatively, invest) in case of default (alternatively, no default), whereas this is not the case under Γ .

When it comes to disclosures in financial markets, Theorem 1 implies that optimal policies should combine the announcement of a pass/fail result (captured by $r \in \{0, 1\}$) with the disclosure of additional information (captured by s) whose role is to guarantee that, when a pass grade is given, the extra information s the agents receive from the policymaker makes investing the unique rationalizable action. This structure appears broadly consistent with common practice. The theorem, however, says more. It indicates that optimal disclosure policies should be transparent about market responses but not in the sense of creating conformism in beliefs about fundamentals. Rather, they should leave no room to ambiguity as to whether or not default will be averted when a pass grade is announced. Preserving heterogeneity in beliefs about fundamentals is key to minimizing the probability of default.

3.2 Pass/fail

Our next result provides a foundation for policies that take a simple pass/fail form; it identifies a key property of the agents' beliefs under which such policies are optimal.

THEOREM 2. *Suppose that $p(x|\theta)$ is log-supermodular. Then, given any regular policy Γ satisfying the perfect-coordination property, there exists a regular binary policy $\Gamma^* = (\{0, 1\}, \pi^*)$ that also satisfies the perfect-coordination property and such that, when agents play according to MARP under both Γ and Γ^* , for any θ , the probability of default and the payoffs (for each agent and the policymaker) are the same under Γ^* and Γ .¹²*

As anticipated in the [Introduction](#), the log-supermodularity of $p(x|\theta)$ (equivalently, the assumption that the distribution $p(x|\theta)$ from which the agents' private signals are drawn satisfies MLRP)) implies that the policymaker cannot reverse the ranking in the agents' optimism through public announcements. Whenever agent j is more optimistic than agent i (in the monotone likelihood-ratio order) based on her exogenous private information x_j , she continues to be more optimistic after hearing the policymaker's announcement, irrespectively of the shape of the policy Γ . In turn, this implies that MARP is always in monotone strategies, and hence that the policymaker does not benefit from disclosing any information beyond the fate of the regime $r^\Gamma(\theta, s)$.

To see this more formally, take any policy $\Gamma = (\mathcal{S}, \pi)$ satisfying the perfect coordination property. Given the result in [Theorem 1](#), without loss of optimality, assume that $\Gamma = (\mathcal{S}, \pi)$ is such that $\mathcal{S} = \{0, 1\} \times S$, for some Polish space S , and that, under MARP consistent with Γ , when the policymaker discloses any signal $(s, r^\Gamma(\theta, s)) = (s, 1)$, investing is the unique rationalizable action for each agent, irrespectively of their exogenous private information. Given the policy Γ , let $U^\Gamma(x, (s, 1)|k)$ denote the expected payoff differential of an agent with exogenous private information x who receives public information $(s, r^\Gamma(\theta, s)) = (s, 1)$ and who expects all other agents to invest if and only if their exogenous signal exceeds a cut-off k . No matter the shape of the policy Γ , when $p(x|\theta)$ is log-supermodular, then MARP associated with the policy Γ is in monotone (i.e., cut-off) strategies. Hence, each agent's expected payoff differential when all other agents play according to MARP can be written as $U^\Gamma(x, (s, 1)|k)$ for some k that depends on s . That the original policy Γ satisfies the perfect-coordination property in turn implies that, for any s and k such that $(k, (s, 1))$ are mutually consistent,¹³ $U^\Gamma(k, (s, 1)|k) > 0$. That is, the expected payoff differential of any agent whose private signal x coincides with the cutoff k must be strictly positive. If this were not the case, the continuation game would also admit a rationalizable profile (in fact, a continuation equilibrium) in which some of the agents refrain from investing, thereby contradicting the fact that investing irrespectively of x is the unique rationalizable profile following the announcement of $(s, 1)$.

Now consider a policy Γ^* that, for any θ , draws the signal $(s, 1)$ (alternatively, $(s, 0)$) from the distribution $\pi(\theta)$ of the original policy $\Gamma = (\mathcal{S}, \pi)$ but conceals the information

¹²The property that $p(x|\theta)$ is log-supermodular means that, for any $x', x'' \in \mathbb{R}$, with $x' < x''$, and any $\theta', \theta'' \in \Theta$, with $\theta'' > \theta'$, then $p(x''|\theta'')p(x'|\theta') \geq p(x''|\theta')p(x'|\theta'')$.

¹³This means that the set $\theta \in \Theta$ such that (a) $k \in \varrho_\theta$ and (b) $(s, 1) \in \text{supp}(\pi(\theta))$ has strictly positive measure under F .

s and only discloses $r = 1$ (alternatively, $r = 0$). By the law of iterated expectations, for all k with $(k, (s, 1))$ mutually consistent, because $U^\Gamma(k, (s, 1)|k) > 0$ then $U^{\Gamma^*}(k, 1|k) > 0$. This implies that the new policy Γ^* also satisfies the perfect-coordination property. The policymaker can thus drop the additional signals s from the original policy Γ and still guarantee that after $r = 1$ is announced, investing is the unique rationalizable action for all agents. That the probability of default and the payoffs (for each agent and the policymaker) are the same under Γ and Γ^* then follows directly from the fact that, for any θ , the probability that each agent invests is the same under the two policies, along with the fact that signals are payoff-irrelevant when fixing the agents' behavior.

The inability to change the ranking in the agents' beliefs through public announcements is key to the optimality of simple pass/fail policies, as the next example shows.

EXAMPLE 1. Suppose that θ is drawn from a uniform distribution over $[-1, 2]$. Given θ , each agent $i \in [0, 1]$ receives an exogenous signal $x_i \in \{x^L, x^H\}$, drawn independently across agents from a Bernoulli distribution with probability

$$p(x^L|\theta) = \begin{cases} 2/3 & \text{if } \theta \in (0, 1/3) \cup [2/3, 5/6) \cup [1, 7/6) \cup [4/3, 5/3) \\ 1/3 & \text{if } \theta \in [1/3, 2/3) \cup [5/6, 1) \cup [7/6, 4/3) \cup [5/3, 2). \end{cases}$$

The value of $p(x^L|\theta)$ for $\theta \in [-1, 0]$ plays no role in this example, so it can be taken arbitrarily. Suppose that agents' payoffs are such that $g(\theta) = 1 - c$ and $b(\theta) = -c$, for all θ , with $c \in (1/2, 8/15)$. There exists a deterministic policy that satisfies PCP and guarantees that default does not occur for $\theta > 0$, whereas no pass/fail policy can guarantee that default does not occur for all $\theta > 0$.¹⁴ $\diamond$

PROOF OF EXAMPLE 1. Figure 1 illustrates the signal structure considered in Example 1. The dash line depicts the probability of signal x^L whereas the solid line the complementary probability of signal x^H , as a function of θ .

Note that the agents' posterior beliefs under the signal structure of Example 1 can be ranked according to FOSD, but not according to MLRP. Each agent observing x^H has posterior beliefs about θ that dominate those of each agent observing x^L in the FOSD order. Nonetheless, the ratio $p(x^H|\theta)/p(x^L|\theta)$ is not increasing in θ over the entire domain, meaning that $p(x|\theta)$ is not log-supermodular and hence posteriors cannot be ranked according to MLRP. Also note that, under the payoff specification in the example, investing is optimal for an agent assigning probability to default no greater than $1 - c$, whereas not investing is optimal if such a probability is at least $1 - c$.

To see that there exists no pass/fail policy guaranteeing that default does not occur for all $\theta > 0$, note that, by virtue of Theorem 1, if such a policy existed, there would also exist a binary policy satisfying PCP and such that $\pi(1|\theta) = 0$ for all $\theta \leq 0$ and $\pi(1|\theta) = 1$ for all $\theta > 0$, with $\pi(1|\theta)$ denoting the probability that the policy discloses signal 1 when

¹⁴The example features signals drawn from a distribution with finite support. This property, however, is not essential. Conclusions similar to those in the example obtain when the agents' signals are drawn from a continuous distribution. We thank Tommaso Denti for suggesting a similar example with finite signals and Leifu Zhang for suggesting an example with continuous signals.

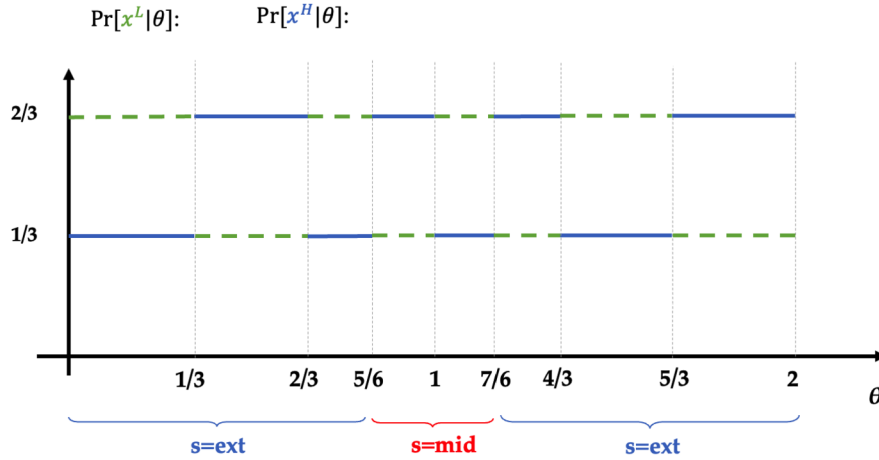


FIGURE 1. Suboptimality of simple pass/tail tests.

the fundamentals are θ . Under such a policy, after hearing that $s = 1$, no matter the private signal x , each agent assigns probability $1/2$ to $\theta \in [0, 1]$ and probability $1/2$ to $\theta \in [1, 2]$. Because $c > 1/2$, each agent expecting all other agents to refrain from investing (and hence default to occur for all $\theta \in [0, 1]$), then finds it optimal to do the same. Hence, under MARP consistent with the above policy, after the signal $s = 1$ is announced, all agents refrain from investing, meaning that the above policy fails to spare types $\theta \in [0, 1]$ from default, when the agents play adversarially.

To see that, instead, the policymaker can avoid default for all $\theta > 0$ using a richer policy, consider the policy $\Gamma = (S, \pi)$, with $S = \{0, (1, \text{mid}), (1, \text{ext})\}$ that, in addition to publicly announcing a pass grade, also announces whether the fundamentals are *extreme* (i.e., $\theta \in (0, 5/6) \cup (7/6, 2]$), or *intermediate* (i.e., $\theta \in [5/6, 7/6]$). Formally, for any $\theta \in [-1, 0]$, $\pi(0|\theta) = 1$, meaning that the policymaker assigns a failing grade. For any $\theta \in [5/6, 7/6]$, instead, $\pi(1, \text{mid}|\theta) = 1$, meaning that the the policymaker announces a pass grade and that fundamentals are intermediate. Finally, for any $\theta \in (0, 5/6) \cup (7/6, 2]$, $\pi(1, \text{ext}|\theta) = 1$, meaning that the policymaker announces a pass grade and that fundamentals are extreme. See Figure 1 for a graphical representation.

Under such a policy, investing is the unique rationalizable action for any agent observing a pass grade, no matter whether the agent also learns that the fundamentals are intermediate or extreme.

To see this, consider first the case in which the fundamentals are extreme, that is, $\theta \in (0, 5/6) \cup (7/6, 2]$. All agents with exogenous information x^H find it dominant to invest when hearing $s = (1, \text{ext})$. In fact, even if all other agents refrained from investing, the probability that each agent with signal x^H assigns to $\theta > 1$ (and hence to the event that there is no default) is $\Pr[\theta > 1|x^H, \text{ext}] = 8/15 > c$, making it dominant to invest. As a consequence of this property, each agent with exogenous private information x^L finds it *iteratively* dominant to invest. This is because, for any $\theta \in [1/3, 5/6]$, even if all agents with exogenous information equal to x^L refrained from investing, the aggregate investment from those individuals with information x^H would suffice for default not to

occur. This means that the probability that each agent with information x^L assigns to the event that default does not occur is at least equal to $\Pr[\theta > 1/3 | x^L, (1, \text{ext})] = 11/15$, implying that it is optimal for the agent to invest.

Next, consider the case in which fundamentals are intermediate, that is, $\theta \in [5/6, 7/6]$. In this case, the ranking of the agents' optimism is reversed, with those agents observing the x^L signal assigning higher probability to higher states. In particular, because each agent with information x^L assigns probability $2/3 > c$ to $\theta \geq 1$, any such agent finds it dominant to invest. Because, for any $\theta \in (5/6, 1)$, $1/3$ of the agents receive information x^L , the minimal size of investment that each agent with signal equal to x^H can expect at any $\theta \in (5/6, 1)$ is equal to $p(x^L | \theta) = 1/3 > 1 - \theta$, implying that even if all the less optimistic agents with signal x^H refrained from investing, default would not occur. But this means that investing is iteratively dominant for those agents receiving the x^H signal.

Hence, the proposed policy spares any $\theta > 0$ from default. Because all agents invest when they observe a pass grade, no matter whether they learn that the fundamentals are extreme or intermediate, one may find it surprising that the policymaker needs to provide the extra information. This is a consequence of the policymaker not trusting the market to play favorably to her. The extra information is precisely what guarantees the uniqueness of the rationalizable action. $\square$

As anticipated above, the benefits from disclosing information in addition to the pass (or fail) grade stem from the possibility to reverse the ranking of the agents' optimism, which is possible only when the distribution $p(x|\theta)$ is not log-supermodular. In the example above, the most optimistic agents are those observing the x^L signals when the fundamentals are intermediate, whereas they are those observing the x^H signals when the fundamentals are extreme. The reversal in the agents' optimism in turn permits the policymaker to guarantee that investing is the unique rationalizable action over a larger set of fundamentals (the entire set $\theta > 0$ in the example).

The above example also illustrates the failure of the Revelation Principle when the policymaker is concerned with unique implementation (equivalently, when the market is expected to play according to MARP). It is well known that, in this case, confining attention to policies that take the form of action recommendations is *with* loss of generality. The contribution of Theorem 2 is in showing that, notwithstanding such a qualification, the optimal policy does take the form of action recommendations in the special case in which beliefs comove with fundamentals according to MLRP.

3.3 Monotone rules

We now turn to the optimality of policies that fail with certainty institutions with weak fundamentals and pass with certainty those with strong fundamentals. As anticipated in the [Introduction](#), the optimality of such rules crucially depends on whether the policymaker's preferences for avoiding default when fundamentals are large are strong enough to compensate for the possibility that nonmonotone rules may permit her to reduce the ex-ante probability of default (i.e., the possibility that default may occur over a set of fundamentals of smaller ex ante probability under a nonmonotone rule).

In this subsection, we identify a condition relating the policymaker's preferences to the agents' exogenous beliefs and payoffs under which monotone rules are optimal. We show that the condition is fairly sharp in that, when violated, one can identify economies in which nonmonotone rules do strictly better than monotone ones. These economies include many of the examples considered in the literature, for example, [Goldstein and Huang \(2016\)](#).

We assume hereafter that

$$\left\{ x \in \mathbb{R} : \int_{\Theta} u(\theta, 1 - P(x|\theta)) \mathbf{1}(\theta > 0) p(x|\theta) dF(\theta) \leq 0 \right\} \neq \emptyset. \quad (1)$$

When Condition (1) is violated, the expected payoff differential between investing and not investing is positive for any agent who is informed that fundamentals are non-negative and who expects each other agent to invest (alternatively, not invest) when receiving a signal above (alternatively, below) hers. In this case, the information-design problem is uninteresting because the policymaker can save all $\theta > 0$ through a policy that announces whether or not $\theta > 0$. Then, let

$$x_{\max} \equiv \sup \left\{ x \in \mathbb{R} : \int_{\Theta} u(\theta, 1 - P(x|\theta)) \mathbf{1}(\theta > 0) p(x|\theta) dF(\theta) \leq 0 \right\}. \quad (2)$$

As we show in the [Appendix](#), $x_{\max}$ is an upper bound for the set of cut-offs characterizing the strategies consistent with MARP across all disclosure policies Γ satisfying the perfect coordination property.

For any x , let $\Theta(x) \equiv \{\theta \in \Theta : x \in \mathcal{Q}_{\theta}\}$ denote the set of fundamentals that, given the distribution $P(\cdot|\theta)$ from which the agents' signals are drawn, are consistent with private information x .

CONDITION M. *The following properties hold:*

- (i) $\inf \Theta(x_{\max}) \leq 0$;
- (ii) for any $\theta_0, \theta_1 \in [0, 1]$, with $\theta_0 < \theta_1$, and $x \leq x_{\max}$ such that (a) $\theta_1 \leq P(x|\theta_1)$ and (b) $x \in \mathcal{Q}_{\theta_0}$,

$$\frac{U^P(\theta_1, 1) - U^P(\theta_1, 0)}{U^P(\theta_0, 1) - U^P(\theta_0, 0)} > \frac{p(x|\theta_1)b(\theta_1)}{p(x|\theta_0)b(\theta_0)}. \quad (3)$$

Property (i) in Condition [M](#) says that the lower bound of the support of the beliefs of an agent with signal $x_{\max}$, where $x_{\max}$ is the threshold defined in (2), is nonpositive and, therefore, that according to this agent there is a positive probability that default is unavoidable, no matter the aggregate investment. Clearly, this property trivially holds when, for any θ , the agents' signals are drawn from a distribution whose support is large enough (and hence, a fortiori, when the noise in the agents' signals is drawn from a distribution with unbounded support, e.g., a Normal distribution).

Property (ii) of Condition [M](#) says that the value the policymaker assigns to avoiding default increases with the underlying fundamentals at a large enough rate. Specifically, the property requires that the benefit that the policymaker derives from changing the

agents' behavior (inducing all agents to invest starting from a situation in which no agent invests) must increase with the fundamentals at a sufficiently high rate, with the critical rate determined by a combination of the agents' payoffs in case of default and beliefs.

THEOREM 3. *Suppose that $p(x|\theta)$ is log-supermodular and Condition **M** holds. Given any regular policy Γ satisfying the perfect-coordination property, there exists a regular deterministic binary monotone policy $\Gamma^{\hat{\theta}} = (\{0, 1\}, \pi^{\hat{\theta}})$ that also satisfies the perfect-coordination property and such that, when the agents play according to MARP under both Γ and $\Gamma^{\hat{\theta}}$, the policymaker's ex ante expected payoff is weakly higher under $\Gamma^{\hat{\theta}}$ than under Γ .¹⁵*

When Condition **M** holds, the choice of the optimal policy reduces to the choice of the smallest threshold $\hat{\theta}$ such that, when agents commonly learn that $\theta > \hat{\theta}$, under the unique rationalizable profile, all agents invest irrespective of their exogenous private information. For this to be the case, it must be that, $\int_{\hat{\theta}}^{\infty} u(\theta, 1 - P(x|\theta))p(x|\theta) dF(\theta) > 0$, for any $x \in \mathbb{R}$.

The above problem, however, does not have a formal solution, due to the lack of upper semicontinuity of the policymaker's payoff in $\hat{\theta}$. Notwithstanding these complications, hereafter we follow the pertinent literature and refer to the “*optimal monotone policy*” as the one defined as follows. For any $\theta \in (0, 1)$, let $x^*(\theta)$ be the critical signal threshold such that, when agents follow a cut-off strategy with threshold $x^*(\theta)$, default occurs if and only if the fundamentals are below θ .¹⁶ Let

$$\theta^* \equiv \inf \left\{ \hat{\theta} \geq 0 : \int_{\hat{\theta}}^{\infty} u(\tilde{\theta}, 1 - P(x^*(\theta)|\tilde{\theta}))p(x^*(\theta)|\tilde{\theta}) dF(\tilde{\theta}) \geq 0 \text{ for all } \theta \in [\hat{\theta}, 1) \right\} \quad (4)$$

be the lowest truncation point $\hat{\theta}$ such that, when the policy reveals that fundamentals are above $\hat{\theta}$, then for any possible default threshold $\theta \in [\hat{\theta}, 1)$, if default were to occur for fundamentals below θ and not for fundamentals above θ , then the marginal agent with signal $x^*(\theta)$ would find it optimal to invest. Hereafter, we assume that θ^* is well-defined, which is always the case when¹⁷

$$\theta^{##} \equiv \sup \left\{ \theta \in (0, 1) : \int_{\theta} u(\tilde{\theta}, 1 - P(x^*(\theta)|\tilde{\theta}))p(x^*(\theta)|\tilde{\theta}) dF(\tilde{\theta}) \leq 0 \right\} < 1.$$

¹⁵The policy $\Gamma^{\hat{\theta}}$ is such that there exists a threshold $\hat{\theta} \in [0, 1]$ such that, for any $\theta \leq \hat{\theta}$, $\pi^{\hat{\theta}}(\theta)$ assigns probability one to $s = 0$, whereas for any $\theta > \hat{\theta}$, $\pi^{\hat{\theta}}(\theta)$ assigns probability one to $s = 1$.

¹⁶For any $\theta \in (0, 1)$, the threshold $x^*(\theta)$ is implicitly defined by $P(x^*(\theta)|\theta) = \theta$. When the noise in the agents' signals is bounded, the definition of $x^*(\theta)$ can be extended to $\theta = 0$ and $\theta = 1$. When the noise is unbounded, abusing notation, one can extend the definition to $\theta = 0$ and $\theta = 1$ by letting $x^*(0) = -\infty$ and $x^*(1) = +\infty$.

¹⁷For any $\hat{\theta} \in (\theta^{##}, 1)$, and any $\theta \in [\hat{\theta}, 1)$,

$$0 < \int_{\theta} u(\tilde{\theta}, 1 - P(x^*(\theta)|\tilde{\theta}))p(x^*(\theta)|\tilde{\theta}) dF(\tilde{\theta}) < \int_{\hat{\theta}}^{\infty} u(\tilde{\theta}, 1 - P(x^*(\theta)|\tilde{\theta}))p(x^*(\theta)|\tilde{\theta}) dF(\tilde{\theta}).$$

Hence, when $\theta^{##} < 1$, θ^* is well-defined.

The *optimal monotone policy* is the one with cut-off $\hat{\theta} = \theta^*$.¹⁸

The previous literature (e.g., Goldstein and Huang (2016)) characterized the threshold θ^* by restricting attention to monotone rules. The contribution of Theorem 3 is in identifying the conditions under which such rules are optimal. Importantly, these conditions are not met in the works that restrict attention to monotone rules. As the examples below suggest, in those settings, the policymaker can strictly increase her payoff through a nonmonotone rule.

As we show in the Appendix, Property (i) in Condition M guarantees that, starting from the optimal monotone policy (the one with cut-off θ^*), one cannot perturb the policy by assigning a pass grade also to a small interval of fundamentals $[\theta', \theta'']$, with $0 \leq \theta' < \theta'' < \theta^*$, while guaranteeing that investing remains the unique rationalizable action when the policymaker announces a pass grade (i.e., when the signal $s = 1$ is disclosed). This property trivially holds when the noise in the agents' signals is large (and hence, a fortiori, when noise is unbounded), but plays a key role when the noise is drawn from a bounded interval of small size (see Example 2 below for an illustration).

Property (ii) of Condition M in turn guarantees that the higher payoff the policymaker obtains, under the new policy, from avoiding default when fundamentals are stronger compensates for the possibility that, from an ex ante perspective, the probability of default may be larger under monotone policies than under nonmonotone ones (see Example 3 for an illustration of why nonmonotone rules may permit the policymaker to avoid default over a set of fundamentals of larger ex ante probability).

As anticipated above, Condition M is fairly sharp in the sense that, when violated, one can identify economies in which the optimal policy is nonmonotone. We provide two such examples below. Example 2 illustrates the role of Property (i) in Condition M, whereas Example 3 illustrates the role of Property (ii) in Condition M. These examples also illustrate why nonmonotone rules, in general, may reduce the set of fundamentals over which default happens.

Let $\theta^{\text{MS}} \in (0, 1)$ be implicitly defined by the unique solution to

$$\int_0^1 u(\theta^{\text{MS}}, A) dA = 0. \quad (5)$$

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 aggregate investment is indifferent between investing and not investing.¹⁹ Importantly, θ^{MS} is independent of the initial common prior F and of the distribution of the agents' signals.

¹⁸The reason why this is an abuse is that, under the monotone policy with cut-off θ^* , in the continuation game that starts after the policymaker announces $s = 1$, there exists a rationalizable profile in which some of the agents refrain from investing. However, there exists a monotone policy with cut-off $\hat{\theta}$ arbitrarily close to the threshold θ^* such that, after the policymaker announces $s = 1$ (equivalently, that $\theta \geq \hat{\theta}$), the unique rationalizable profile features all agents investing. Because the policymaker's payoff under the latter policy is arbitrarily close to the one she obtains when all agents invest for $\theta > \theta^*$ and refrain from investing when $\theta \leq \theta^*$, the abuse appears justified.

¹⁹This means that the agent believes that aggregate investment is uniformly distributed over $[0, 1]$. See Morris and Shin (2006).

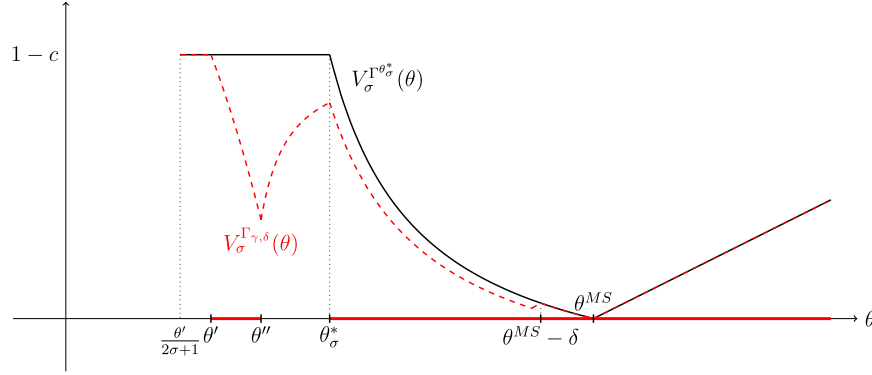


FIGURE 2. Suboptimality of deterministic binary monotone policies.

EXAMPLE 2. Suppose that there exist scalars $g, b \in \mathbb{R}$, with $g > 0 > b$, such that, for any θ , $g(\theta) = g$, and $b(\theta) = b$. Assume that θ is drawn from a uniform distribution with support $[-K, 1 + K]$, for some $K \in \mathbb{R}_{++}$. Finally, assume that the agents' exogenous signals are given by $x_i = \theta + \sigma \epsilon_i$, with $\sigma \in \mathbb{R}_{++}$ and with each ϵ_i drawn independently across agents from a uniform distribution over $[-1, 1]$, with $\sigma < K/2$. Let θ_{σ}^* be the threshold defined in (4), applied to the primitives described in this example.²⁰ There exists $\sigma^{\#} \in (0, K/2)$ such that (a) $\inf \Theta(x_{\sigma^{\#}}^*(\theta^{MS})) > 0$, and (b) for all $\sigma \in (0, \sigma^{\#})$, starting from the optimal monotone policy with cut-off θ_{σ}^* , there exists a deterministic nonmonotone policy satisfying the perfect-coordination property and permitting the policymaker to avoid default over a set of fundamentals of strictly larger probability measure than the optimal monotone policy. $\diamond$

The proof is in the Supplementary Appendix. Here, we sketch the key arguments. To fix ideas, let $g = 1 - c$ and $b = -c$, with $c \in (0, 1)$, as in Example 1, and recall that, under such a payoff specification, investing is optimal when the probability of default is no greater than $1 - c$, whereas not investing is optimal when such a probability exceeds $1 - c$.

For any binary policy $\Gamma = (\{0, 1\}, \pi)$, and any threshold $\theta \in [0, 1]$ such that $(x_{\sigma}^*(\theta), 1)$ are mutually consistent under Γ , let

$$V_{\sigma}^{\Gamma}(\theta) \equiv U_{\sigma}^{\Gamma}(x_{\sigma}^*(\theta), 1 | x_{\sigma}^*(\theta)),$$

denote the payoff of the marginal agent with signal $x_{\sigma}^*(\theta)$, after the policy Γ announces that $s = 1$, where U_{σ}^{Γ} is the function defined after Theorem 2.

Now, for any $\hat{\theta} \in \Theta$, let $\Gamma^{\hat{\theta}} = (\{0, 1\}, \pi^{\hat{\theta}})$ be the deterministic, binary, monotone rule with cut-off $\hat{\theta}$. Note that the absence of any public disclosure is equivalent to a monotone policy with cut-off $\hat{\theta} = \min \Theta = -K$ and that, under such a policy, default occurs if and only if $\theta \leq \theta^{MS} = c$.

²⁰Hereafter, the subscript σ in θ_{σ}^* and x_{σ}^* is meant to highlight that these thresholds are those for the economy in which the noise in the agents' exogenous private signals is scaled by σ .

A necessary and sufficient condition for all agents to invest under MARP consistent with the policy $\Gamma^{\hat{\theta}}$, after hearing that $s = 1$, is that, for any possible default threshold $\theta > \hat{\theta}$, $V_{\sigma}^{\Gamma^{\hat{\theta}}}(\theta) > 0$. The lowest fundamental in the support of $x_{\sigma}^*(\theta)$'s beliefs is $x_{\sigma}^*(\theta) - \sigma$. Hence, when $x_{\sigma}^*(\theta) - \sigma > \hat{\theta}$, the marginal agent with signal $x_{\sigma}^*(\theta)$ *already knows* from his private information that fundamentals are above $\hat{\theta}$. Because, in the absence of any public disclosure, the payoff of the marginal agent is strictly negative for all $\theta < \theta^{\text{MS}}$, this implies that the cut-off θ_{σ}^* for the optimal monotone rule is $\theta_{\sigma}^* = x_{\sigma}^*(\theta^{\text{MS}}) - \sigma$.

Now to see that the optimal monotone policy is improvable, assume that σ is small so that $x_{\sigma}^*(\theta^{\text{MS}}) - \sigma > 0$. Next, pick $\gamma, \delta > 0$ small and let $\theta'' \equiv x_{\sigma}^*(\theta^{\text{MS}} - \delta) - \sigma$ and $\theta' \equiv \theta'' - \gamma$, with $\theta' > 0$. Consider a binary policy $\Gamma_{\gamma, \delta} = (\{0, 1\}, \pi_{\gamma, \delta})$, that in addition to announcing a pass grade $s = 1$ when fundamentals are above θ_{σ}^* (as the optimal monotone rule does) also announces $s = 1$ when $\theta \in [\theta', \theta'']$. Let $V_{\sigma}^{\Gamma_{\gamma, \delta}}(\theta)$ be the payoff of the marginal agent with signal $x_{\sigma}^*(\theta)$ under the new rule $\Gamma_{\gamma, \delta}$, after the policymaker announces that $s = 1$. This payoff is represented in Figure 2 along with the payoff $V_{\sigma}^{\Gamma^{\theta_{\sigma}^*}}(\theta)$ under the optimal monotone rule. Provided that γ and δ are small, $V_{\sigma}^{\Gamma_{\gamma, \delta}}(\theta) \geq 0$ for all θ for which $(x_{\sigma}^*(\theta), 1)$ are mutually consistent under $\Gamma_{\gamma, \delta}$, with $V_{\sigma}^{\Gamma_{\gamma, \delta}}(\theta) = 0$ if and only if $\theta = \theta^{\text{MS}}$. Starting from $\Gamma_{\gamma, \delta}$, one can then further perturb the policy $\Gamma_{\gamma, \delta}$ by giving a fail grade to banks with fundamentals in $[\theta_{\sigma}^*, \theta_{\sigma}^* + \varepsilon]$, with $\varepsilon > 0$ small. The new policy $\tilde{\Gamma}$ so constructed is such that $V_{\sigma}^{\tilde{\Gamma}}(\theta) > 0$ for all θ for which $(x_{\sigma}^*(\theta), 1)$ are mutually consistent under $\tilde{\Gamma}$, meaning that, when the policymaker announces that $s = 1$, investing is the *unique* rationalizable action for all agents. The policy $\tilde{\Gamma}$ thus satisfies the perfect-coordination property and guarantees that default occurs over a set of fundamentals of strictly smaller probability under F than the optimal monotone policy $\Gamma^{\theta_{\sigma}^*} = (\{0, 1\}, \pi^{\theta_{\sigma}^*})$.

The reason why the nonmonotone policy $\tilde{\Gamma}$ constructed in the proof of Example 2 guarantees that default occurs over a smaller set of fundamentals than the optimal monotone policy is that agents receiving signals around θ^{MS} are highly sensitive to the grade the policy gives to institutions with fundamentals around θ^{MS} but not so much so to the grade given to fundamentals far from θ^{MS} . In the above example with bounded noise, an agent receiving a signal $x_{\sigma}^*(\theta^{\text{MS}})$ is not sensitive at all to the grade the policy gives to fundamentals below $x_{\sigma}^*(\theta^{\text{MS}}) - \sigma$ because his private signal informs him that the fundamentals are above $x_{\sigma}^*(\theta^{\text{MS}}) - \sigma$. Hence, while it is impossible to amend the optimal monotone policy (the one with cut-off $\theta_{\sigma}^* = x_{\sigma}^*(\theta^{\text{MS}}) - \sigma$) by giving a pass grade also to fundamentals slightly below θ_{σ}^* without inducing some of the agents to refrain from investing, it is possible to amend the optimal monotone policy by extending the pass grade to an interval $[\theta', \theta'']$ of fundamentals sufficiently “far away” from θ_{σ}^* , while continuing to induce all agents to invest under MARP. The reason why such improvements are not feasible under Condition M in Theorem 3 is that Property (i) in Condition M implies that $x_{\sigma}^*(\theta^{\text{MS}}) - \sigma < 0$, thus making the above construction unfeasible.²¹ Interestingly, when $\theta \in [\theta', \theta'']$, the assumption of bounded support of the agents' beliefs

²¹Under Property (i), the marginal agent with signal $x_{\sigma}^*(\theta^{\text{MS}})$ does not rule out any fundamental in $(0, \theta^{\text{MS}})$. Hence, any perturbation of the optimal monotone policy passing fundamentals to the left of θ^{MS} induces the agent to refrain from investing.

implies that a positive-measure set of agents know with certainty that $\theta \in [\theta', \theta'']$ and yet, under the unique rationalizable profile, all agents invest; this is because, by design, the policy $\tilde{\Gamma}$ constructed in Example 2 guarantees that, when $\theta \in [\theta', \theta'']$, such an event is not commonly learned.

The next example considers an economy in which the noise in the agents' exogenous signals is drawn from a distribution with an unbounded support (in which case, Property (i) in Condition M trivially holds), but Property (ii) is violated.

Given any binary, deterministic policy $\Gamma = (\{0, 1\}, \pi)$ (i.e., any policy such that, for any θ , $\pi(\theta)$ is a degenerate Dirac distribution assigning probability 1 either to $s = 1$ or to $s = 0$), let $D^\Gamma = \{(\underline{\theta}_i, \bar{\theta}_i] : i = 1, \dots, N\}$ denote the partition of $(0, \theta^{\text{MS}}]$ induced by π , with $N \in \mathbb{N}$, $\underline{\theta}_1 = 0$, and $\bar{\theta}_N = \theta^{\text{MS}}$.²² Let $d \in D^\Gamma$ denote a generic cell of the partition D^Γ and, for any $\theta \in (0, \theta^{\text{MS}}]$, denote by $d^\Gamma(\theta) \in D^\Gamma$ the cell that contains θ . Finally, let $M(\Gamma) \equiv \max_{i=1, \dots, N} |\bar{\theta}_i - \underline{\theta}_i|$ denote the *mesh* of D^Γ , that is, the Lebesgue measure of the cell of D^Γ of maximal Lebesgue measure.

Example 3 below shows that, when the noise in the agents' information is small, any deterministic binary policy of large mesh can be improved upon by a non-monotone deterministic binary policy with a smaller mesh. This property in turn implies that optimal policies are highly nonmonotone.

EXAMPLE 3. Suppose that θ is drawn from an improper uniform prior over $\mathbb{R}$ and that the agents' signals are given by $x_i = \theta + \sigma \varepsilon_i$, with ε_i drawn from a standard Normal distribution.²³ Further assume that there exist scalars $g, b, W, L \in \mathbb{R}$, with $g > 0 > b$ and $W > L$, such that, for any θ , $g(\theta) = g$, $b(\theta) = b$, $W(\theta) = W$ and $L(\theta) = L$. There exists a scalar $\bar{\sigma} > 0$ and a function $\mathcal{E} : (0, \bar{\sigma}] \rightarrow \mathbb{R}_+$, with $\lim_{\sigma \rightarrow 0^+} \mathcal{E}(\sigma) = 0$, such that, for any $\sigma \in (0, \bar{\sigma}]$, in the game in which the noise in the agents' information is scaled by σ , the following is true: given any deterministic binary policy $\Gamma = (\{0, 1\}, \pi)$ satisfying the perfect-coordination property and such that $M(\Gamma) > \mathcal{E}(\sigma)$, there exists another deterministic binary policy Γ^* with $M(\Gamma^*) < \mathcal{E}(\sigma)$ that also satisfies the perfect-coordination property and such that the ex ante probability of default under Γ^* is strictly smaller than under Γ . $\diamond$

See the Supplementary Appendix for a detailed proof of the result. Here, we discuss the main ideas. Nonmonotone policies permit the policymaker to avoid default over a larger set of fundamentals by making it difficult for the agents to commonly learn the fundamentals when the latter are between 0 and θ^{MS} and the policymaker announces a pass grade. Intuitively, if the policymaker assigned a pass grade to an interval $(\theta', \theta''] \subset (0, \theta^{\text{MS}}]$ of large Lebesgue measure, when σ is small and $\theta \in (\theta', \theta'']$, most agents would receive private signals $x_i \in (\theta', \theta'']$. No matter the grade assigned to fundamentals outside the interval $(\theta', \theta'']$, in the continuation game that starts after the policymaker announces a pass grade, most agents with signals $x_i \in (\theta', \theta'']$ would then

²²That is, either (a) $\pi(\theta) = 0$ for all $\theta \in \cup_{i=2k, k \leq N} (\underline{\theta}_i, \bar{\theta}_i]$ and $\pi(\theta) = 1$ for all $\theta \in \cup_{i=2k-1, k \leq N} (\underline{\theta}_i, \bar{\theta}_i]$, or (b) $\pi(\theta) = 1$ for all $\theta \in \cup_{i=2k, k \leq N} (\underline{\theta}_i, \bar{\theta}_i]$ and $\pi(\theta) = 0$ for all $\theta \in \cup_{i=2k-1, k \leq N} (\underline{\theta}_i, \bar{\theta}_i]$.

²³The improperness of the prior simplifies the exposition but is not important. The agents' hierarchies of beliefs are still well-defined.

assign high probability to the joint event that $\theta \in (\theta', \theta'']$, that other agents assign high probability to $\theta \in (\theta', \theta'']$, and so on. When this is the case, it is rationalizable for such agents to refrain from investing. Hence, when σ is small, the only way the policymaker can guarantee that, when $\theta \in (0, \theta^{\text{MS}}]$, the agents invest after hearing a pass grade is by dividing the set $(0, \theta^{\text{MS}}]$ into a collection of disjoint intervals, each of small Lebesgue measure. This guarantees that the support of each agent's posterior beliefs after a pass grade is announced is not connected. Connectedness of the supports facilitates rationalizable profiles where some agents refrain from investing.

Next, suppose that the intervals $(\underline{\theta}_i, \bar{\theta}_i] \subset (0, \theta^{\text{MS}}]$, $i = 1, \dots, N$, receiving a pass grade are far apart, implying that the policymaker fails an interval $(\theta', \theta''] \subset (0, \theta^{\text{MS}}]$ of large Lebesgue measure (note that this is indeed the case under the optimal monotone deterministic rule with cutoff θ_σ^* , where θ_σ^* is the threshold defined in (4)).²⁴ The detailed derivations in the Supplementary Appendix then show that, starting from Γ , the policymaker could assign a pass grade to fundamentals in the middle of $(\theta', \theta'']$ and a fail grade to some fundamentals to the right of θ'' , in such a way that (a) investing continues to be the unique rationalizable action for all agents after hearing a pass grade, and (b) the set of fundamentals receiving a pass grade under the new policy is strictly larger than under the original one. Furthermore, the construction sketched above can be iterated until one arrives at a new policy with a mesh smaller than $\mathcal{E}(\sigma)$ under which default occurs over a set of fundamentals of strictly smaller measure than under the original policy. When the benefit $W(\theta) - L(\theta)$ of avoiding default is constant in θ , as in the example above, the new policy thus yields the policymaker a strictly higher payoff than the original one.

Finally, one can show that, when σ is small, a pass grade can be given to all $\theta > \theta^{\text{MS}} + \varepsilon$, with $\varepsilon > 0$ small, while guaranteeing that all agents invest after the policymaker announces the pass grade $s = 1$.²⁵

The above properties thus also imply that, if the policymaker is restricted to deterministic policies (arguably, the most relevant case in practice), when the precision of the agents' exogenous information is large, the optimal policy is highly nonmonotone over $(0, \theta^{\text{MS}})$ and announces a pass grade when fundamentals are above θ^{MS} .

4. EXTENSIONS

We first introduce a few enrichments in Section 4.1, then establish the analog of the three theorems above for these richer economies in Section 4.2, and then conclude in Section 4.3 discussing the role of the multiplicity of the receivers and their exogenous private information.

²⁴The subscript simply highlights the dependence of the cutoff θ_σ^* on σ .

²⁵Formally, for any $\varepsilon > 0$, there exists $\sigma(\varepsilon)$ such that, for any $\sigma < \sigma(\varepsilon)$, given any pass/fail policy Γ satisfying PCP, there exists another pass/fail policy Γ' also satisfying PCP that agrees with Γ on any $\theta < \theta^{\text{MS}}$ and gives a pass grade to any $\theta \geq \theta^{\text{MS}} + \varepsilon$.

4.1 Generalizations

The fundamentals are given by (θ, z) , with θ drawn from Θ according to the absolutely continuous cdf F , and with z drawn from $[\underline{z}, \bar{z}]$ according to $Q_\theta(z)$, with the cdf $Q_\theta(z)$ weakly decreasing in θ , for any z .²⁶

The variable θ continues to parameterize the maximal information the policymaker can collect about the fundamentals. The additional variable z parameterizes risk that the agents and the policymaker face at the time of the disclosure (e.g., macroeconomic variables that are only imperfectly correlated with the fundamentals). As in the baseline model, conditional on θ , the private signals $\mathbf{x} = (x_i)_{i \in [0,1]}$ are i.i.d. draws from an (absolutely continuous) cumulative distribution function $P(x|\theta)$, with associated density $p(x|\theta)$ strictly positive and bounded over the interval $Q_\theta \in \mathbb{R}$.

There exists a function $R : \Theta \times [0, 1] \times [\underline{z}, \bar{z}] \rightarrow \mathbb{R}$ such that, given any (θ, A, z) , default occurs (i.e., $r = 0$) if, and only if, $R(\theta, A, z) \leq 0$. The function R is continuous and strictly increasing in (θ, z, A) . For any (θ, A) , the probability of avoiding default is thus given by $r(\theta, A) \equiv \mathbb{P}[R(\theta, A, z) > 0 | \theta, A]$.

There exist functions $\hat{W}, \hat{L} : \Theta \times [0, 1] \times [\underline{z}, \bar{z}] \rightarrow \mathbb{R}$ such that, given any (θ, A, z) , the policymaker's payoff is equal to

$$\hat{U}^P(\theta, A, z) = \hat{W}(\theta, A, z)\mathbf{1}(R(\theta, A, z) > 0) + \hat{L}(\theta, A, z)\mathbf{1}(R(\theta, A, z) \leq 0). \quad (6)$$

Hence, $\hat{W}(\theta, A, z)$ is the policymaker's payoff in case default is avoided, whereas $\hat{L}(\theta, A, z)$ is her payoff in case of default. Likewise, there exist functions $\hat{g}, \hat{b} : \Theta \times [0, 1] \times [\underline{z}, \bar{z}] \rightarrow \mathbb{R}$ such that, given any (θ, A, z) , the agents' payoff differential between investing and not investing is equal to

$$\hat{u}(\theta, A, z) = \hat{g}(\theta, A, z)\mathbf{1}(R(\theta, A, z) > 0) + \hat{b}(\theta, A, z)\mathbf{1}(R(\theta, A, z) \leq 0), \quad (7)$$

with $\hat{g}(\theta, A, z) > 0 > \hat{b}(\theta, A, z)$, for any (θ, A, z) . For any (θ, A) , then let

$$g(\theta, A) \equiv \frac{\mathbb{E}[\mathbf{1}(R(\theta, A, z) > 0)\hat{g}(\theta, A, z)|\theta, A]}{r(\theta, A)} \quad \text{and} \\ b(\theta, A) \equiv \frac{\mathbb{E}[\mathbf{1}(R(\theta, A, z) \leq 0)\hat{b}(\theta, A, z)|\theta, A]}{1 - r(\theta, A)}$$

denote the agents' expected payoff differential in case of no default and in case of default, respectively. Likewise, for any (θ, A) , let

$$W(\theta, A) \equiv \frac{\mathbb{E}[\mathbf{1}(R(\theta, A, z) > 0)\hat{W}(\theta, A, z)|\theta, A]}{r(\theta, A)} \quad \text{and} \\ L(\theta, A) \equiv \frac{\mathbb{E}[\mathbf{1}(R(\theta, A, z) \leq 0)\hat{L}(\theta, A, z)|\theta, A]}{1 - r(\theta, A)}$$

²⁶All the results extend to the case where $Q_\theta(z)$ has unbounded support. Note that $Q_\theta(z)$ is not required to be absolutely continuous in z (in fact, it is not absolutely continuous in the baseline model, where the distribution has a mass point of 1 at $z = 0$).

denote the policymaker's expected payoff, again in case of no default and default, respectively.

The agents' and the policymaker's expected payoffs can then be conveniently expressed as a function of θ and A only, by letting

$$u(\theta, A) \equiv r(\theta, A)g(\theta, A) + (1 - r(\theta, A))b(\theta, A) \quad \text{and} \\ U^P(\theta, A) \equiv r(\theta, A)W(\theta, A) + (1 - r(\theta, A))L(\theta, A).$$

Hereafter, we assume that $|u(\theta, A)|$ is bounded and that there exist $\underline{\theta}, \bar{\theta} \in \mathbb{R}$, with $\underline{\theta} < \bar{\theta}$, such that (a) $u(\theta, 1) < 0$ for all $\theta \leq \underline{\theta}$, (b) $u(\theta, 0) > 0$ for all $\theta > \bar{\theta}$, and (c) $u(\theta, 1) > 0 > u(\theta, 0)$ for all $\theta \in (\underline{\theta}, \bar{\theta}]$. The thresholds $\underline{\theta}$ and $\bar{\theta}$ define the “critical region” $(\underline{\theta}, \bar{\theta}]$ where the sign of the agents' payoff differential depends on the response of the market.²⁷ We also assume that both $u(\theta, A)$ and $U^P(\theta, A)$ are nondecreasing in A and such that $U^P(\theta, 1) > U^P(\theta, 0)$ for all $\theta \in (\underline{\theta}, \bar{\theta}]$.²⁸

4.2 Results

We identify conditions under which Theorems 1–3 extend to these richer economies.

4.2.1 Perfect-coordination property Given any distribution $G \in \Delta\Theta$ over Θ , say that G is “regular” if, when the common posterior over Θ is G and, for any θ , agents receive private signals according to $P(\cdot|\theta)$, MARP is well-defined. Then, for any regular G , any θ , let $A(\theta; G)$ denote the aggregate investment at θ when agents play according to MARP, under the common posterior G .

CONDITION PC. For any distribution $\tau \in \Delta\Delta(\Theta)$ over posterior beliefs consistent with the common prior F (i.e., such that $\int G\tau(dG) = F$), the following condition holds:

$$\int \left(\int [\mathbf{1}(u(\theta, A(\theta; G)) > 0)U^P(\theta, 1) + \mathbf{1}(u(\theta, A(\theta; G)) \leq 0)U^P(\theta, 0)]G(d\theta) \right) \tau(dG) \\ \geq \int \left(\int U^P(\theta, A(\theta; G))G(d\theta) \right) \tau(dG).$$

²⁷The critical region can also be defined in terms of the regime outcome. That is, let $\underline{\theta}', \bar{\theta}' \in \mathbb{R}$, with $\underline{\theta}' < \bar{\theta}'$, be defined by $R(\underline{\theta}', 1, \bar{z}) = R(\bar{\theta}', 0, \underline{z}) = 0$. Note that default occurs with certainty when $\theta < \underline{\theta}'$ and never occurs when $\theta > \bar{\theta}'$, no matter (A, z) . Because the agents' payoff differential is strictly negative (alternatively, strictly positive) when there is default (alternatively, when there is no default), $(\underline{\theta}, \bar{\theta}] \subseteq (\underline{\theta}', \bar{\theta}']$. All the results below hold also under this alternative definition. The reason for defining the critical region as done above is that it permits us to weaken some of the assumptions by requiring that they hold over a smaller set of fundamentals. Clearly, the two definitions coincide when the regime outcome is a deterministic function of (θ, A) , as in the baseline model.

²⁸That $u(\theta, A)$ is monotone in A implies that the continuation game remains supermodular. That $U^P(\theta, A)$ is nondecreasing in A implies that, for any Γ , MARP continues to coincide with the “smallest” rationalizable profile, that is, the one involving the smallest measure of agents investing. Finally, that for any θ in the critical region, the policymaker strictly prefers that all agents invest to no agent investing guarantees that, when the optimal policy has a pass/fail structure, it is obtained by maximizing the probability that a pass grade is given when fundamentals are in the critical range.

To appreciate the meaning of Condition [PC](#), suppose that the policymaker, through her disclosure policy Γ , generates a distribution τ over common posteriors G over Θ , and that, for any G , agents play according to MARP. Now suppose that, in each state θ , the policymaker also informs the agents of the sign of their expected payoff differential $u(\theta, A(\theta; G))$ under MARP consistent with G . Finally, suppose that, after each posterior G is generated, the additional information induces all agents to invest when they learn that $u(\theta, A(\theta; G)) > 0$ and not to invest when they learn that $u(\theta, A(\theta; G)) \leq 0$. Then the additional information makes the agents better off. Condition [PC](#) says that the policymaker is also weakly better off. In other words, the condition requires that the policymaker's and the agents' payoffs be not too misaligned. Condition [PC](#) trivially holds when the policymaker faces no aggregate uncertainty (i.e., when each distribution Q_θ over $[\underline{z}, \bar{z}]$ is degenerate), W is weakly increasing in A and L is invariant in A , as in the baseline model. For example, in case of stress testing, the condition says that the policymaker prefers more agents to invest in case the bank under examination avoids default, but is indifferent as to how many investors pull their money out of the bank when the latter defaults. More generally, Condition [PC](#) accommodates for the possibility that both W and L depend on A , possibly nonmonotonically, provided that, on average, the loss to the policymaker from having no agent invest in states θ in which the agents' expected payoff differential (under MARP given the induced common posterior G) is negative is more than compensated by the benefit from having all agents invest in states θ in which the differential is positive. The average is over both the induced posteriors G and the fundamentals θ .

As in the baseline model, let $A^\Gamma(\theta, s)$ denote the aggregate size of investment at θ under MARP consistent with Γ , when the policy discloses s .

THEOREM 1*. *Given any regular policy $\Gamma = (S, \pi)$, there exists another regular policy Γ^* satisfying the perfect-coordination property and such that when under both Γ and Γ^* agents play according to MARP the following are true: (1) for any θ , no agent is worse off under Γ^* than under Γ , and some agents are strictly better off; (2) if, for any θ and $s \in \text{supp}(\pi(\theta))$, the regime outcome is deterministic (i.e., $r(\theta, A^\Gamma(\theta, s)) \in \{0, 1\}$), then for any θ the probability of default under Γ^* is the same as under Γ ; (3) when Condition [PC](#) holds, the policymaker is better off under Γ^* than under Γ .*

Theorem 1* extends Theorem 1 to the richer class of economies under consideration, in which the regime outcome is determined by additional variables that are not observable by the policymaker, and where both the policymaker's and the agents' payoffs depend on the aggregate investment A beyond its effect on the regime outcome. The policy Γ^* in the theorem is obtained from the original policy Γ by disclosing, for each θ , in addition to the information $s \in \text{supp}(\pi(\theta))$ disclosed by the original policy Γ , a second piece of information that reveals to the market whether at (θ, s) , under MARP consistent with the original policy Γ , the agents' expected payoff differential is positive or negative. Note in particular that because the sign of the payoff differential in the baseline model is given by the regime outcome, this additional piece of information, in the baseline model, coincides with the regime outcome $r^\Gamma(\theta, s) \in \{0, 1\}$.

In [Inostroza and Pavan \(2024a\)](#), we show that the perfect-coordination property is fairly general and extends to a class of economies even richer than the one introduced in Section 4.1 in which (a) the agents' prior beliefs need not be consistent with a common prior, nor be generated by signals drawn independently across agents, conditionally on θ , (b) the number of agents is arbitrary (in particular, finitely many agents), (c) payoffs can be heterogenous across agents, (d) agents have a level-K degree of sophistication, (e) the policymaker may possess imperfect information about the payoff state and/or the agents' beliefs, (f) the policymaker may engage in flexible discriminatory disclosures and disclose different information to different agents. The key property is the possibility for the policymaker to have access to information that is a *sufficient statistic* of the agents' information when predicting the sign of the agents' payoff differential under MARP. This property holds when, for example, the correlation in the agents' exogenous beliefs originates in public signals the policy maker has access to.²⁹

4.2.2 Pass/fail policies

CONDITION FB. *For any x , $u(\theta, 1 - P(x|\theta)) \geq 0$ (alternatively, $u(\theta, 1 - P(x|\theta)) \leq 0$) implies that $u(\theta', 1 - P(x|\theta')) > 0$ for all $\theta' > \theta$ (alternatively, $u(\theta', 1 - P(x|\theta')) < 0$ for all $\theta' < \theta$).*

Condition **FB** (which stands for “single crossing from below”) states that, for any x , the payoff differential $u(\theta, 1 - P(x|\theta))$ from investing when all agents follow a cut-off strategy with cut-off x crosses 0 once *from below*. The property clearly holds in the baseline model where (i) $r(\theta, A) = \mathbf{1}(A > 1 - \theta)$ and (ii) $g(\theta) > 0 > b(\theta)$ for all θ . It also holds when $u(\theta, A)$, in addition to being nondecreasing in A as assumed above, is nondecreasing in θ .

THEOREM 2*. *Suppose that $p(x|\theta)$ is log-supermodular and Condition **FB** holds. Then, given any regular policy $\Gamma = (\mathcal{S}, \pi)$ satisfying the perfect-coordination property, there exists a regular binary policy $\Gamma^* = (\{0, 1\}, \pi^*)$ that also satisfies the perfect-coordination property and such that, when agents play according to MARP under both Γ and Γ^* , for any θ , the probability of default and the payoffs (for each agent and the policymaker) are the same under Γ^* and Γ .*

Because $\Gamma = (\mathcal{S}, \pi)$ satisfies the perfect-coordination property, $\cup_{\theta} \text{supp}(\pi(\theta))$ can be partitioned in two sets, $\mathcal{S}_1$ and $\mathcal{S}_0$, such that, under Γ , all agents invest (alternatively, do not invest) when receiving information $s \in \mathcal{S}_1$ (alternatively, $s \in \mathcal{S}_0$), irrespectively of their private signals x . The key step in the proof in the [Appendix](#) shows that the log-supermodularity of $p(x|\theta)$, together with *Condition FB*, jointly imply that, under *any* policy, MARP is in cut-off strategies. The reason is the same as the one discussed above for the baseline model. In turn, this property implies that all agents continue to invest

²⁹We conjecture that, as long as the above sufficient statistic property holds, Theorem 2* and 3* below also extend to settings in which the agents' signals are not conditionally independent given θ . Whether the results extend to some environments in which the sufficient statistic property is violated is an interesting question for future work.

(alternatively, refrain from investing) when the policymaker “pools the signals” and discloses only that $s \in S_1$ (alternatively, that $s \in S_0$). The arguments are similar to those leading to Theorem 2 above. The policy $\Gamma^* = (\{0, 1\}, \pi^*)$ is then constructed by letting $\pi^*(1|\theta) = \pi(S_1|\theta)$ (and $\pi^*(0|\theta) = \pi(S_0|\theta)$) for all θ . Contrary to the baseline model, after the policy Γ^* discloses signal 1 (alternatively, signal 0), the regime outcome need not be deterministic. Nonetheless, the probability of default is the same under the two policies Γ and Γ^* and so are the payoffs.³⁰

4.2.3 Monotone rules First, we extend the definition of $x_{\max}$ to accommodate for the fact that, in richer economies, $\underline{\theta}$ need not coincide with 0. That is, we let

$$x_{\max} \equiv \sup \left\{ x \in \mathbb{R} : \int_{\Theta} u(\theta, 1 - P(x|\theta)) \mathbf{1}(\theta > \underline{\theta}) p(x|\theta) dF(\theta) \leq 0 \right\}. \quad (8)$$

Next, we extend Condition M as follows.

CONDITION M*. (i*) $\inf \Theta(x_{\max}) \leq \underline{\theta}$;

(ii*) For any $\theta_0, \theta_1 \in [\underline{\theta}, \bar{\theta}]$, with $\theta_0 < \theta_1$, and $x \leq x_{\max}$ such that (a) $u(\theta_1, 1 - P(x|\theta_1)) \leq 0$ and (b) $x \in \Theta_{\theta_0}$,

$$\frac{U^P(\theta_1, 1) - U^P(\theta_1, 0)}{U^P(\theta_0, 1) - U^P(\theta_0, 0)} > \frac{p(x|\theta_1)u(\theta_1, 1 - P(x|\theta_1))}{p(x|\theta_0)u(\theta_0, 1 - P(x|\theta_0))}. \quad (9)$$

(iii*) $|u(\theta, 1 - P(x|\theta))|$ is log-supermodular over $\{(\theta, x) \in [\underline{\theta}, \bar{\theta}] \times \mathbb{R} : u(\theta, 1 - P(x|\theta)) \leq 0\}$.³¹

Property (i*) is similar to Property (i) in Condition M in the baseline model but accommodates for the fact that, in richer economies, $\underline{\theta}$ need not coincide with 0. Property (ii*) extends Property (ii) in Condition M to the current environment with richer preferences in which $u(\theta, A)$ and $U^P(\theta, A)$ depend on A over and above the effect that the latter variable has on the regime outcome.

Property (iii*) is a new condition that requires that, for any $\theta' < \theta''$ and $x' < x''$ such that $u(\theta'', 1 - P(x'|\theta'')) < 0$,

$$\frac{u(\theta'', 1 - P(x'|\theta''))}{u(\theta', 1 - P(x'|\theta'))} \leq \frac{u(\theta'', 1 - P(x''|\theta''))}{u(\theta', 1 - P(x''|\theta'))}. \quad (10)$$

Note that $u(\theta'', 1 - P(x'|\theta'')) < 0$ implies that $u(\theta', 1 - P(x'|\theta'))$, $u(\theta', 1 - P(x''|\theta'))$, $u(\theta'', 1 - P(x''|\theta'')) < 0$. The condition thus requires that the relative reduction in the expected losses stemming from the fundamentals improving from θ' to $\theta'' > \theta'$ is larger

³⁰As in the baseline model, that payoffs (for each agent and the policymaker) are the same under Γ and Γ^* follows from the fact that, for any θ , the probability that each agent invests is the same under the two policies, along with the fact that signals are payoff-irrelevant when fixing the agents' behavior.

³¹The log-supermodularity of $|u(\theta, 1 - P(x|\theta))|$ means that, for any $x', x'' \in \mathbb{R}$, with $x' < x''$, and any $\theta', \theta'' \in \Theta$, with $\theta'' > \theta'$, such that $u(\theta'', 1 - P(x'|\theta'')) \leq 0$,

$$u(\theta'', 1 - P(x''|\theta''))u(\theta', 1 - P(x'|\theta')) \geq u(\theta'', 1 - P(x'|\theta''))u(\theta', 1 - P(x''|\theta')).$$

when agents invest if and only if $x > x'$ than when they invest if and only if $x > x'' > x'$. The reduction in the losses $u(\theta, 1 - P(x|\theta))$ combines the direct effect of θ on $u(\theta, A)$ with the indirect effect of θ on $A = 1 - P(x|\theta)$ that obtains when the agents follow a cut-off strategy whereby they invest if and only if their signals exceed x . The condition trivially holds in the baseline model where $u(\theta, A) < 0$ if and only if, given (θ, A) , there is default (i.e., $A \leq 1 - \theta$), in which case $u(\theta, A) = b(\theta)$.

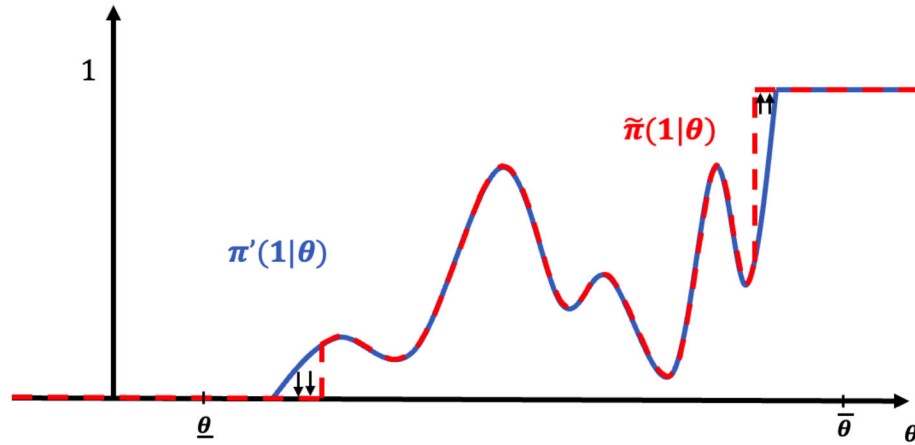
THEOREM 3*. *Suppose that $p(x|\theta)$ is log-supermodular and Conditions PC, FB, and M^* hold. Given any regular policy Γ , there exists a regular deterministic binary monotone policy $\Gamma^{\hat{\theta}} = (\{0, 1\}, \pi^{\hat{\theta}})$ that satisfies the perfect-coordination property and such that, when the agents play according to MARP under both Γ and $\Gamma^{\hat{\theta}}$, the policymaker's ex ante expected payoff is weakly higher under $\Gamma^{\hat{\theta}}$ than under Γ .*

To gain some intuition on the role played by the additional requirement in Condition M^* (property (iii*)), first observe that the conditions in the theorem imply that Theorems 1* and 2* hold. Given any policy Γ , there thus exists a binary policy $\Gamma' = (\{0, 1\}, \pi')$ satisfying the perfect-coordination property and such that $\pi'(1|\theta) = 0$ for all $\theta \leq \underline{\theta}$ and $\pi'(1|\theta) = 1$ for all $\theta > \bar{\theta}$ and such that the policymaker is weakly better off under Γ' than under Γ . The policy Γ' can be constructed following the steps in the proofs of Theorems 1* and 2*. Now suppose that Γ' is not a deterministic monotone rule (i.e., there is no $\hat{\theta}$ such that $\pi'(1|\theta) = \mathbf{1}(\theta \geq \hat{\theta})$ for F -almost all θ). As in Section 3.2, let $U^{\Gamma'}(x, 1|x)$ be the expected payoff of an agent with signal x who, under the policy Γ' hears that $s = 1$, and who expects all other agents to invest if and only if their signal exceeds x . Suppose that $U^{\Gamma'}(x, 1|x)$ has a unique global minimum $\bar{x} \equiv \arg \min_x U^{\Gamma'}(x, 1|x)$, and that $\bar{x} \leq x_{\max}$ (these properties are not assumed in the proof but permit us to illustrate the role of Property (iii*) in Condition M^* in the simplest possible terms). That Γ' satisfies the perfect-coordination property implies that $U^{\Gamma'}(\bar{x}, 1|\bar{x}) > 0$, for otherwise it is rationalizable for some of the agents with signal $x \leq \bar{x}$ not to invest, after hearing that $s = 1$ (the arguments are similar to those in the baseline model). To make things interesting, suppose there exist two disjoint intervals of fundamentals $\Theta^-, \Theta^+ \subset \Theta(\bar{x})$, both consistent with $\bar{x}$, such that (a) $\sup \Theta^- \leq \inf \Theta^+$, (b) $u(\theta, 1 - P(\bar{x}|\theta)) \leq 0$ for F -almost all $\theta \in \Theta^- \cup \Theta^+$, (c) $\pi'(1|\theta) > 0$ for F -almost all $\theta \in \Theta^-$, and (iv) $\pi'(1|\theta) < 1$ for F -almost all $\theta \in \Theta^+$ (as we show in the Appendix, when these properties do not hold there exist trivial improvements of the policy Γ' even when Property (iii*) in Condition M^* does not hold).

Then, consider a binary policy $\tilde{\Gamma} = (\{0, 1\}, \tilde{\pi})$ constructed from Γ' by reducing the probability of the pass grade $s = 1$ over the interval Θ^- and increasing it over the interval Θ^+ , as in Figure 3. Let

$$\Delta S(x) \equiv \int_{-\infty}^{+\infty} u(\theta, 1 - P(x|\theta)) p(x|\theta) (\tilde{\pi}(1|\theta) - \pi'(1|\theta)) dF(\theta).$$

Suppose that the new policy $\tilde{\Gamma}$ is such that $\Delta S(\bar{x}) = 0$, which implies that $U^{\tilde{\Gamma}}(\bar{x}, 1|\bar{x}) > 0$. Property (iii*), along with the fact that (a) $\tilde{\pi}(1|\theta) - \pi'(1|\theta)$ crosses zero from below, (b) $p(x|\theta)$ is log-supermodular, and (c) Conditions FB holds, guarantees that $\Delta S(x) \geq 0$ for all $x < \bar{x}$, which in turn implies that $U^{\tilde{\Gamma}}(x, 1|x) > 0$ for all $x < \bar{x}$. The proof in the

FIGURE 3. Construction of improving policy $\tilde{\Gamma} = (\{0, 1\}, \tilde{\pi})$.

Appendix leverages this property to show how to construct a sequence of perturbations of the policy Γ' leading to a new binary policy $\Gamma^{\hat{\theta}} = (\{0, 1\}, \pi^{\hat{\theta}})$ that is deterministic and monotone and such that $U^{\Gamma^{\hat{\theta}}}(x, 1|x) > 0$ for all x , which guarantees that $\Gamma^{\hat{\theta}}$ also satisfies the perfect-coordination property. That the new policy $\Gamma^{\hat{\theta}}$ improves over the original one Γ then follows from the fact that the policymaker's payoff satisfies Property (ii*)—the arguments for this last step are similar to those leading to Theorem 3 in the baseline model.

4.3 Discussion: Role of multiplicity of receivers and exogenous private information

It is worth contrasting the results about the suboptimality of monotone rules (when Condition M^* is violated) to those for economies featuring either a single privately-informed receiver, or multiple receivers with no exogenous private information.

Single receiver. With a single receiver, the optimal policy is a simple monotone pass/fail policy with cutoff equal to $\theta^* = 0$. This is because, in this model, the policymaker's and the receiver's payoffs are aligned (they both want to avoid default when possible). With a single receiver, there is no risk of adversarial coordination, and hence the optimal policy coincides with the one that the designer would select if she trusted the receiver to play favorably to her.

Things are different when preferences are misaligned. To see this, suppose the policymaker's payoff is equal to W in case of no default, and $L < W$ in case of default, with $W, L \in \mathbb{R}$ constant, as in Examples 2 and 3 above. However, now suppose that the receiver's payoff differential between investing and not investing is equal to $-g$ in case of default and $-b$ in case of no default, with $g > 0 > b$. Such a payoff differential may reflect the idea that the receiver is a speculator whose payoff is zero when he refrains from speculating (equivalently, when he invests), is positive when he speculates and default occurs, and is negative when he speculates and default does not occur. Using the results in Guo and Shmaya (2019), one can then show that the optimal policy in this

case has the *interval structure*: each type x of the receiver is induced to play the action favorable to the policymaker (abstain from speculating) over an interval of fundamentals $[\theta_1(x), \theta_2(x)]$, with $\theta_1(x) < 1 < \theta_2(x)$, for all x , and with $\theta_1(x)$ decreasing in x and $\theta_2(x)$ increasing in x . Such a policy requires disclosing more than two signals, and hence cannot be implemented through a simple pass/fail test. In contrast, with a continuum of heterogeneously informed receivers with the same payoffs as in the variant above, the optimal policy is a pass–fail test that is typically nonmonotone in θ .³² Furthermore, when the optimal policy is not monotone, it does not have the interval structure, as each receiver with signal x is induced to invest over a nonconnected set of fundamentals. The reason for these differences is that, with a single receiver, to discourage the latter from taking the adversarial action, the policymaker must persuade the receiver that the fundamentals are likely to be above 1, in which case the attack is unsuccessful. With multiple receivers, instead, the policymaker must persuade each receiver that enough other receivers are not attacking, which as shown above, is best accomplished by a nonmonotone policy that makes it difficult for the receivers to commonly learn the fundamentals, when the latter are between 0 and θ^{MS} .³³

Multiple receivers with no exogenous private information. When all receivers have the same posterior beliefs, no matter whether payoffs are aligned or misaligned, under MARP, each receiver plays the friendly action only if it is dominant to do so. The optimal policy is a simple monotone pass/fail policy with cutoff θ^* implicitly defined by

$$\int_{\theta^*}^1 b \, dF(\theta) + \int_1^{\infty} g \, dF(\theta) = 0.$$

The reason why the optimal policy is monotone when the receivers possess no exogenous private information is that the policymaker needs to convince each of them that θ is above 1 with sufficiently high probability to make the friendly action dominant.

5. CONCLUSIONS

We consider the design of public information in coordination settings in which the designer does not trust the receivers to play favorably to her. We show that, despite the fear of adversarial coordination, the optimal policy induces all receivers to take the same action. Importantly, while each agent can perfectly predict the action of any other agent, he is not able to predict the beliefs that rationalize such actions. We identify conditions under which the optimal policy has a pass/fail structure, as well as conditions under which the optimal policy is monotone, passing institutions with strong fundamentals and failing the others.

³²This is because, under MARP, all agents play the friendly action if and only if it is iteratively dominant for them to do so, irrespective of the alignment in payoffs.

³³Mensch (2021) characterizes general conditions under which the optimal policy is monotone with a single, *uninformed*, receiver. Goldstein and Leitner (2018) study an economy in which these conditions are not satisfied and the optimal policy is nonmonotone. The analysis in these works is very different in that it does not identify the role that coordination and the receivers' private information play for the optimal policy.

The results are worth extending in a few directions. The analysis assumes that the policymaker is Bayesian and knows the distribution from which the agents' exogenous private information is drawn. While this is a natural starting point, in future work it would be interesting to investigate how the structure of the optimal policy is affected by the policymaker's uncertainty about the agents' information sources.³⁴

Motivated by the applications the analysis is meant for (most notably, stress testing), we have confined attention to nondiscriminatory disclosures. In future work, it would be interesting to extend the analysis to settings in which agents are endowed with exogenous private information (as assumed here) but the designer can disclose different information to different agents (discriminatory policies).

The analysis in the present paper is static. Many applications of interest are dynamic, with agents coordinating on multiple attacks and/or learning over time (for the role of dynamics in global games, see, among others, [Angeletos, Hellwig, and Pavan \(2007\)](#)). In future work, it would be interesting to consider dynamic extensions and investigate how the timing of information disclosures is affected by the agents' behavior in previous periods.³⁵

Finally, the analysis is conducted by assuming that the maximal information that the designer can collect about the fundamentals (in the paper, θ) is exogenous. In future work, it would be interesting to accommodate for the possibility that part of this information is endogenous. For example, in stress testing, the policymaker may solicit information from the same banks that are under scrutiny. This creates an interesting screening + persuasion problem in the spirit of the literature on privacy in sequential contacting (see, e.g., [Calzolari and Pavan \(2006a,b\)](#), and [Dworczak \(2020\)](#)).³⁶

APPENDIX

PROOF OF THEOREM 1*. Given any regular policy $\Gamma = (S, \pi)$ and any $n \in \mathbb{N}$, let $T_{(n)}^\Gamma$ be the set of strategies surviving n rounds of iterated deletion of interim strictly dominated strategies (IDISDS), with $T_{(0)}^\Gamma$ denoting the entire set of strategy profiles $\mathbf{a} = (a_i(\cdot))_{i \in [0, 1]}$, where for any $i \in [0, 1]$, $a_i(x, s)$ denotes the probability agent i invests, given (x, s) . Let $\mathbf{a}_{(n)}^\Gamma \equiv (a_{(n),i}^\Gamma(\cdot))_{i \in [0, 1]} \in T_{(n)}^\Gamma$ denote the most aggressive profile surviving n rounds of IDISDS (i.e., the profile in $T_{(n)}^\Gamma$ that is most adversarial to the policymaker, in the sense that it minimizes the policymaker's ex ante payoff). The profiles $(\mathbf{a}_{(n)}^\Gamma)_{n \in \mathbb{N}}$ can be constructed inductively as follows. The profile $\mathbf{a}_{(0)}^\Gamma \equiv (a_{(0),i}^\Gamma(\cdot))_{i \in [0, 1]}$ prescribes that all

³⁴See [Dworczak and Pavan \(2022\)](#) for a notion of robustness in information design that accounts for this type of ambiguity.

³⁵For models of dynamic persuasion, see, among others, [Ely \(2017\)](#) and [Basak and Zhou \(2022\)](#).

³⁶[Calzolari and Pavan \(2006a\)](#) considers an auction setting in which the sender is the initial owner of a good and where the different receivers are privately-informed bidders in an upstream market who then resell in a downstream market. [Calzolari and Pavan \(2006b\)](#) studies information design in a model of sequential contracting with multiple principals, where upstream principals play the role of senders persuading downstream principals (the receivers). [Dworczak \(2020\)](#) contains a general analysis of persuasion in mechanism-design environments with aftermarkets in which senders restrict attention to cut-off mechanisms.

agents refrain from investing, irrespective of (x, s) . Next, let $U_i^\Gamma(x, s; \mathbf{a})$ denote the payoff differential between investing and not investing for agent i receiving information (x, s) when, under Γ , all other agents follow the strategy in $\mathbf{a}$. Then $a_{(n),i}^\Gamma(x, s) = 0$ if $U_i^\Gamma(x, s; \mathbf{a}_{(n-1)}^\Gamma) \leq 0$ and $a_{(n),i}^\Gamma(x, s) = 1$ if $U_i^\Gamma(x, s; \mathbf{a}_{(n-1)}^\Gamma) > 0$. MARP consistent with Γ is the profile $(a_i^\Gamma(\cdot))_{i \in [0,1]}$ given by $a_i^\Gamma(\cdot) = \lim_{n \rightarrow \infty} a_{(n),i}^\Gamma(\cdot)$, all $i \in [0, 1]$.

Next, observe that, for any n , there exists a function $a_{(n),i}^\Gamma(\cdot)$ such that $a_{(n),i}^\Gamma(\cdot) = a_{(n)}^\Gamma(\cdot)$ for all $i \in [0, 1]$. With an abuse of notation, hereafter we thus denote by a^Γ the common strategy that all agents follow under MARP consistent with Γ . For any θ and $s \in \text{supp}(\pi(\theta))$, aggregate investment under MARP consistent with Γ given (θ, s) is thus the same for any $\mathbf{x}, \mathbf{x}' \in \mathbf{X}(\theta)$ and is given by $A^\Gamma(\theta, s) \equiv \int a^\Gamma(x, s) dP(x|\theta)$.

Next, consider the policy $\Gamma^+ = (S^+, \pi^+)$, $S^+ \equiv S \times \{0, 1\}$, that, for each θ , draws the public signal s from the same distribution $\pi(\theta) \in \Delta(S)$ as the original policy Γ , and then, for each s it draws, it also announces the sign of the agents' payoff differential at (θ, s) , when agents play according to MARP consistent with the original policy Γ . That is, for any θ and any $s \in \text{supp}(\pi(\theta))$, the new policy Γ^+ thus announces $(s, \mathbf{1}(u(\theta, A^\Gamma(\theta, s)) > 0))$. In the baseline model of Section 2, the sign of $u(\theta, A^\Gamma(\theta, s))$ is uniquely determined by the regime outcome $r^\Gamma(\theta, s)$. In that environment, for any θ , and any $s \in \text{supp}(\pi(\theta))$, the new policy Γ^+ thus announces $(s, r^\Gamma(\theta, s))$.

Define $T_{(n)}^{\Gamma^+}$ and $a_{(n)}^{\Gamma^+}$ analogously to $T_{(n)}^\Gamma$ and $a_{(n)}^\Gamma$ above, but with respect to Γ^+ .

The proof is in three steps. Steps 1 and 2 show that any agent i who, given (x, s) , finds it dominant (alternatively, iteratively dominant) to invest under Γ , also finds it dominant (alternatively, iteratively dominant) to invest under Γ^+ when receiving information $(x, (s, 1))$. Step 3 uses the above property to establish that, because the game is supermodular and $\mathbf{a}^{\Gamma^+}$ is "less aggressive" than $\mathbf{a}^\Gamma$ (meaning that any agent who, given (x, s) , invests under $\mathbf{a}^\Gamma$ also invests under $\mathbf{a}^{\Gamma^+}$ when receiving information $(x, (s, 1))$), then, under $\mathbf{a}^{\Gamma^+}$, all agents invest (alternatively, refrain from investing) when receiving information $(s, 1)$ (alternatively, $(s, 0)$).

Step 1. We prove that $\{(x, s) : U_i^\Gamma(x, s; \mathbf{a}) > 0 \forall \mathbf{a}\} \subseteq \{(x, s) : U_i^{\Gamma^+}(x, (s, 1); \mathbf{a}) > 0 \forall \mathbf{a}\}$, for all $i \in [0, 1]$. That is, any agent i who, under Γ , finds it dominant to invest, given information (x, s) , also finds it dominant to invest under Γ^+ when receiving information $(x, (s, 1))$.

First, note that the supermodularity of the game implies that $\{(x, s) : U_i^\Gamma(x, s; \mathbf{a}) > 0 \forall \mathbf{a}\} = \{(x, s) : U_i^\Gamma(x, s; \mathbf{a}_{(0)}^\Gamma) > 0\}$ and $\{(x, s) : U_i^{\Gamma^+}(x, (s, 1); \mathbf{a}) > 0 \forall \mathbf{a}\} = \{(x, s) : U_i^{\Gamma^+}(x, (s, 1); \mathbf{a}_{(0)}^{\Gamma^+}) > 0\}$.

Now let $\Lambda_i^\Gamma(x, s)$ denote the distribution over Θ describing the beliefs of agent $i \in [0, 1]$ when receiving information $(x, s) \in \mathbb{R} \times S$ under Γ , and $\Lambda_i^{\Gamma^+}(x, (s, 1))$ the corresponding beliefs under Γ^+ , when receiving information $(x, (s, 1))$ under Γ^+ . Bayesian updating implies that

$$\Lambda_i^{\Gamma^+}(d\theta|x, (s, 1)) = \frac{\mathbf{1}(u(\theta, A^\Gamma(\theta, s)) > 0)}{\Lambda_i^\Gamma(1|x, s)} \Lambda_i^\Gamma(d\theta|x, s), \quad (11)$$

where $\Lambda_i^\Gamma(1|x, s) \equiv \int_{\{\theta \in \Theta : u(\theta, A^\Gamma(\theta, s)) > 0\}} \Lambda_i^\Gamma(d\theta|x, s)$ is the total probability an agent with information (x, s) assigns, under Γ , to fundamentals for which $u(\theta, A^\Gamma(\theta, s)) > 0$.

Next, observe that, for any $i \in [0, 1]$ and $(x, s) \in \mathbb{R} \times \mathcal{S}$ such that

$$U_i^\Gamma(x, s; \mathbf{a}_{(0)}^\Gamma) = \int_\theta u(\theta, 0) \Lambda_i^\Gamma(d\theta|x, s) > 0, \quad (12)$$

we also have that

$$\begin{aligned} U_i^{\Gamma^+}(x, (s, 1); \mathbf{a}_{(0)}^{\Gamma^+}) \Lambda_i^\Gamma(1|x, s) &= \int_\theta u(\theta, 0) \mathbf{1}(u(\theta, A^\Gamma(\theta, s)) > 0) \Lambda_i^\Gamma(d\theta|x, s) \\ &\geq \int_\theta u(\theta, 0) \Lambda_i^\Gamma(d\theta|x_i, s) = U_i^\Gamma(x, s; \mathbf{a}_{(0)}^\Gamma) > 0. \end{aligned}$$

The first equality follows from the fact that, under $\mathbf{a}_{(0)}^\Gamma$, no agent invests, along with the property of posterior beliefs in (11). The first inequality follows from the monotonicity of $u(\theta, A)$ in A along with the fact $A^\Gamma(\theta, s) \geq 0$, which together imply that $u(\theta, 0) \leq 0$ for any θ for which $u(\theta, A^\Gamma(\theta, s)) \leq 0$. The second equality follows from the definition of $U_i^\Gamma(x, s; \mathbf{a}_{(0)}^\Gamma)$. Finally, the second inequality follows from (12).

Thus, any agent for whom investing was dominant after receiving information (x, s) under Γ , continues to find it dominant to invest after receiving information $(x, (s, 1))$ under Γ^+ .

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

$$\{(x, s) : U_i^\Gamma(x, s; \mathbf{a}) > 0 \ \forall \mathbf{a} \in T_{(k-1)}^\Gamma\} \subseteq \{(x, s) : U_i^{\Gamma^+}(x, (s, 1); \mathbf{a}) > 0, \ \forall \mathbf{a} \in T_{(k-1)}^{\Gamma^+}\}. \quad (13)$$

Arguments similar to those establishing the result in Step 1 above imply that

$$\{(x, s) : U_i^\Gamma(x, s; \mathbf{a}) > 0 \ \forall \mathbf{a} \in T_{(n-1)}^\Gamma\} \subseteq \{(x, s) : U_i^{\Gamma^+}(x, (s, 1); \mathbf{a}) > 0, \ \forall \mathbf{a} \in T_{(n-1)}^{\Gamma^+}\}. \quad (14)$$

Intuitively, the result follows from the following three properties: (a) because the game is supermodular, $\{(x, s) : U_i^\Gamma(x, s; \mathbf{a}) > 0 \ \forall \mathbf{a} \in T_{(n-1)}^\Gamma\} = \{(x, s) : U_i^\Gamma(x, s; \mathbf{a}_{(n-1)}^\Gamma) > 0\}$, where recall that $\mathbf{a}_{(n-1)}^\Gamma$ is the most aggressive profile surviving $n-1$ rounds of IDISDS (clearly, the same property holds for Γ^+); (b) $\mathbf{a}_{(n-1)}^{\Gamma^+}$ is “less aggressive” than $\mathbf{a}_{(n-1)}^\Gamma$, in the sense that any agent who, given (x, s) , invests under $\mathbf{a}_{(n-1)}^\Gamma$ also invests under Γ^+ when receiving information $(x, (s, 1))$; and (c) the extra information that θ is such $u(\theta, A^\Gamma(\theta, s)) > 0$ removes from the support of the agents’ posterior beliefs states in which the payoff differential from investing is nonpositive under $\mathbf{a}^\Gamma$, and hence also under $\mathbf{a}_{(n-1)}^\Gamma$ (recall that $\mathbf{a}_{(n-1)}^\Gamma$ is more aggressive than $\mathbf{a}^\Gamma$, meaning that any agent who, given (x, s) , invests under $\mathbf{a}_{(n-1)}^\Gamma$, also invests under $\mathbf{a}^\Gamma$ when receiving the same information (x, s)).

Step 3. Equipped with the results in Steps 1 and 2 above, we now prove that, for all $\theta \in \Theta$ and $s \in \text{supp}(\pi(\theta))$ such that $u(\theta, A^\Gamma(\theta, s)) > 0$, $a^{\Gamma^+}(x, (s, 1)) \equiv \lim_{n \rightarrow \infty} a_{(n)}^{\Gamma^+}(x, (s, 1)) = 1$ for all x . This follows directly from the fact that, as shown above, $a^\Gamma(x, s) = 1 \Rightarrow a^{\Gamma^+}(x, (s, 1)) = 1$. The announcement that θ is such that $u(\theta, A^\Gamma(\theta, s)) > 0$ thus reveals to each agent that, when all other agents play according to MARP consistent with the new policy Γ^+ , the payoff differential from investing is strictly positive. Any agent i receiving information $(s, 1)$ under Γ^+ thus necessarily invests, no matter x . Under the new policy Γ^+ , all agents thus invest when they learn that θ is such that $u(\theta, A^\Gamma(\theta, s)) > 0$.

That they all refrain from investing when they learn that θ is such that $u(\theta, A^\Gamma(\theta, s)) \leq 0$ follows from the fact that such an announcement makes it common certainty that $\theta \leq \bar{\theta}$.

We conclude that the new policy Γ^+ satisfies the perfect-coordination property. That, when the agents play according to MARP, for any θ , no agent is worse off (and some agents are strictly better off) under Γ^+ than under Γ follows from the fact that, for all $s \in \text{supp}(\pi(\theta))$, the following are true: (1) when (θ, s) is such that $u(\theta, A^\Gamma(\theta, s)) > 0$, all agents who are not investing under Γ (thus obtaining an expected payoff of zero) invest under Γ^+ (obtaining an expected payoff $u(\theta, 1) > 0$), and all agents who are investing under Γ continue to invest but obtain a larger payoff $u(\theta, 1) > u(\theta, A^\Gamma(\theta, s))$ because of the monotonicity of $u(\theta, A)$ in A ; (2) when, instead, (θ, s) is such that $u(\theta, A^\Gamma(\theta, s)) \leq 0$, all agents who are not investing under Γ (thus obtaining an expected payoff of zero) continue not to invest under Γ^+ , whereas all agents who are investing under Γ (obtaining a negative payoff) now refrain from investing thus obtaining a payoff of zero.

Next, suppose that, under MARP consistent with Γ , for any θ and $s \in \text{supp}(\pi(\theta))$, the regime outcome is a deterministic function of (θ, s) . Then, for any (θ, s) , the sign of $u(\theta, A^\Gamma(\theta, s))$ is determined by the regime outcome (it is strictly positive when $r^\Gamma(\theta, s) = 1$, i.e., when there is no default, and it is weakly negative when $r^\Gamma(\theta, s) = 0$, i.e., when there is default). Because the regime outcome is monotone in A , by inducing all agents to invest when $u(\theta, A^\Gamma(\theta, s)) > 0$ and not to invest when $u(\theta, A^\Gamma(\theta, s)) \leq 0$, the policy Γ^+ induces the same regime outcome as Γ .

To see that the policymaker is better off under Γ^+ than under Γ , for any set of signals $S \subseteq \mathcal{S}$, any θ , let $\pi^+(S, 1|\theta)$ (alternatively, $\pi^+(S, 0|\theta)$) denote the probability that the policy π^+ selects signals $(s, 1)$ (alternatively, $(s, 0)$) with $s \in S$. Then let $\Pi^{\Gamma^+}(S, 1) \equiv \int_\theta \pi^+(S, 1|\theta) dF(\theta)$ (alternatively, $\Pi^{\Gamma^+}(S, 0) \equiv \int_\theta \pi^+(S, 0|\theta) dF(\theta)$) denote the ex ante probability of announcements $(s, 1)$ (alternatively, $(s, 0)$) with $s \in S$, under the policy Γ^+ . Finally, for any $S \subseteq \mathcal{S}$, let $\Pi^\Gamma(S) \equiv \int \pi(S|\theta) dF(\theta)$ denote the ex ante probability the policy Γ selects signals in S . Condition PC implies that

$$\begin{aligned} & \int_S \left(\int [\mathbf{1}(u(\theta, A^\Gamma(\theta, s)) > 0) U^P(\theta, 1) + \mathbf{1}(u(\theta, A^\Gamma(\theta, s)) \leq 0) U^P(\theta, 0)] \Lambda^\Gamma(d\theta|s) \right) \Pi^\Gamma(ds) \\ & \geq \int_S \left(\int U^P(\theta, A^\Gamma(\theta, s)) \Lambda^\Gamma(d\theta|s) \right) \Pi^\Gamma(ds). \end{aligned}$$

Hence, the policymaker is better off under Γ^+ than under Γ .

The result in the theorem then follows by taking $\Gamma^* = \Gamma^+$. $\square$

PROOF OF THEOREM 2*. The proof is in 2 steps. Step 1 shows that, when $p(x|\theta)$ is log-supermodular and Condition FB holds, then under *any* regular policy, MARP is in cut-off strategies. Step 2 then leverages the result in Step 1 to show that, starting from any policy Γ that satisfies the perfect-coordination property, one can construct a binary policy Γ^* that also satisfies the perfect-coordination property and such that, for any θ , the probability that each agent invests under Γ^* is the same as under Γ , which implies the result in the theorem.

Step 1. Fix an arbitrary policy $\Gamma = (S, \pi)$ and, for any pair $(x, s) \in \mathbb{R} \times S$, let $\Lambda^\Gamma(\theta|x, s)$ represent the endogenous posterior beliefs over Θ of each agent receiving exogenous information x and endogenous information s . Next, let $U^\Gamma(x, s|k) \equiv \int u(\theta, 1 - P(k|\theta))\Lambda^\Gamma(d\theta|x, s)$ denote the expected payoff differential of an agent with information (x, s) , when all other agents follow a cut-off strategy with cut-off k (i.e., they invest if their private signal exceeds k and refrain from investing if it is below k).

LEMMA 1. *Suppose that $p(x|\theta)$ is log-supermodular and that Condition FB holds. Given any policy $\Gamma = (S, \pi)$, for any $s \in S$, there exists $\xi^{\Gamma;s} \in \mathbb{R}$ such that MARP consistent with Γ is given by the strategy profile $\mathbf{a}^\Gamma \equiv (a_i^\Gamma)_{i \in [0,1]}$ such that, for any $s \in S$, $x \in \mathbb{R}$, $i \in [0, 1]$, $a_i^\Gamma(x, s) = \mathbf{1}(x > \xi^{\Gamma;s})$ with $\xi^{\Gamma;s} \equiv \sup\{x : U^\Gamma(x, s|x) \leq 0\}$ if $\{x : U^\Gamma(x, s|x) \leq 0\} \neq \emptyset$, and $\xi^{\Gamma;s} \equiv -\infty$ otherwise. Moreover, the strategy profile $\mathbf{a}^\Gamma$ is a BNE of the continuation game that starts with the announcement of the policy Γ .*

PROOF OF LEMMA 1. Fix the policy $\Gamma = (S, \pi)$. For any $s \in S$, let $\xi_{(1)}^{\Gamma;s} \equiv \sup\{x : \lim_{k \rightarrow \infty} U^\Gamma(x, s|k) \leq 0\}$. Given the public signal s , it is dominant for any agent with private signal x exceeding $\xi_{(1)}^{\Gamma;s}$ to invest. Next, recall that, for any $n \in \mathbb{N}$, $T_{(n)}^\Gamma$ denotes the set of strategy profiles that survive the first n rounds of iterated deletion of interim strictly dominated strategies (IDISDS), and $\mathbf{a}_{(n)}^\Gamma \equiv (a_{(n),i}^\Gamma)_{i \in [0,1]}$ the most aggressive profile in $T_{(n)}^\Gamma$. Observe that the profile $\mathbf{a}_{(1)}^\Gamma$ is given by $a_{(1),i}^\Gamma(x, s) = \mathbf{1}(x > \xi_{(1)}^{\Gamma;s})$ for all $(x, s) \in \mathbb{R} \times S$, and all $i \in [0, 1]$, and minimizes the policymaker's payoff not just in expectation but for any (θ, s) . This follows from the fact that, when nobody else invests, the expected payoff differential $\int u(\theta, 0)\Lambda^\Gamma(d\theta|x, s)$ between investing and not investing crosses 0 only once and from below at $x = \xi_{(1)}^{\Gamma;s}$. The single-crossing property of $\int u(\theta, 0)\Lambda^\Gamma(d\theta|x, s)$ in turn is a consequence of the fact that $u(\theta, 0)$ crosses 0 only once from below at $\theta = \bar{\theta}$ (as implied by Condition FB and the definition of $\bar{\theta}$) along with Property SCB below.

PROPERTY SCB. *Suppose that the function $h : \mathbb{R} \rightarrow \mathbb{R}$ crosses 0 only once from below at $\theta = \theta_0$ (i.e., $h(\theta) \leq 0$ for all $\theta \leq \theta_0$ and $h(\theta) \geq 0$ for all $\theta > \theta_0$). Let $q : \mathbb{R}^2 \rightarrow \mathbb{R}_+$ be a log-supermodular function and suppose that, for any θ , there is an open interval $\varrho_\theta = (\underline{\varrho}_\theta, \bar{\varrho}_\theta) \subset \mathbb{R}$ containing θ such that $q(x, \theta) > 0$ for all $x \in \varrho_\theta$ and $q(x, \theta) = 0$ for (almost) all $x \in \mathbb{R} \setminus \varrho_\theta$, with the bounds $\underline{\varrho}_\theta, \bar{\varrho}_\theta$ nondecreasing in θ . Choose any (Lebesgue) measurable subset $\Omega \subseteq \mathbb{R}$ containing θ_0 and, for any $x \in \mathbb{R}$, let $\Psi(x; \Omega) \equiv \int_\Omega h(\theta)q(x, \theta) d\theta$. Suppose there exists $x^* \in \varrho_{\theta_0}$ such that $\Psi(x^*; \Omega) = 0$. Then, necessarily, $\Psi(x; \Omega) \geq 0$ for all $x \in \varrho_{\theta_0}$ with $x > x^*$, and $\Psi(x; \Omega) \leq 0$ for all $x \in \varrho_{\theta_0}$ with $x < x^*$, with both inequalities strict if (a) $\{\theta \in \Omega : h(\theta) \neq 0\}$ has strict positive Lebesgue measure, (b) q is strictly log-supermodular over $\mathbb{R}^2$.³⁷*

PROOF OF PROPERTY SCB. For any $x \in \mathbb{R}$, let $\Omega_x \equiv \{\theta \in \Omega : x \in \varrho_\theta\}$. The monotonicity of ϱ_θ in θ implies that Ω_x is monotone in x in the strong-order sense. Pick any $x' \in \varrho_{\theta_0}$ with $x' > x^*$. That x^* and x' belong to ϱ_{θ_0} implies that $\theta_0 \in \Omega_{x^*} \cap \Omega_{x'}$. Next, observe that

$$\Psi(x'; \Omega) = \int_{\Omega_{x'}} h(\theta)q(x', \theta) d\theta$$

³⁷That q is strictly log-supermodular over $\mathbb{R}^2$ also implies that $q(x, \theta) > 0$ for all $(x, \theta) \in \mathbb{R}^2$.

$$\begin{aligned}
&= \int_{\Omega_{x'} \cap \Omega_{x^*}} h(\theta) q(x', \theta) d\theta + \int_{\Omega_{x'} \setminus \Omega_{x^*}} h(\theta) q(x', \theta) d\theta \\
&= \int_{\Omega_{x^*} \cap \Omega_{x'} \cap (-\infty, \theta_0)} h(\theta) q(x^*, \theta) \frac{q(x', \theta)}{q(x^*, \theta)} d\theta \\
&\quad + \int_{\Omega_{x^*} \cap \Omega_{x'} \cap (\theta_0, \infty)} h(\theta) q(x^*, \theta) \frac{q(x', \theta)}{q(x^*, \theta)} d\theta \\
&\quad + \int_{\Omega_{x'} \setminus \Omega_{x^*}} h(\theta) q(x', \theta) d\theta \\
&\geq \frac{q(x', \theta_0)}{q(x^*, \theta_0)} \left(\int_{\Omega_{x^*} \cap \Omega_{x'} \cap (-\infty, \theta_0)} h(\theta) q(x^*, \theta) d\theta + \int_{\Omega_{x^*} \cap \Omega_{x'} \cap (\theta_0, \infty)} h(\theta) q(x^*, \theta) d\theta \right) \\
&\quad + \int_{\Omega_{x'} \setminus \Omega_{x^*}} h(\theta) q(x', \theta) d\theta \\
&\geq \frac{q(x', \theta_0)}{q(x^*, \theta_0)} \underbrace{\Psi(x^*; \Omega)}_{=0} + \int_{\Omega_{x'} \setminus \Omega_{x^*}} h(\theta) q(x', \theta) d\theta \geq 0.
\end{aligned}$$

The first equality follows from the fact that $q(x', \theta) = 0$ for almost all $\theta \in \Omega \setminus \Omega_{x'}$. The second equality follows from the fact that $\Omega_{x'}$ can be partitioned into $\Omega_{x'} \cap \Omega_{x^*}$ and $\Omega_{x'} \setminus \Omega_{x^*}$. The third equality follows from noting that $q(x^*, \theta) > 0$ for all $\theta \in \Omega_{x^*}$. The first inequality follows from the monotonicity of $q(x', \theta)/q(x^*, \theta)$ over $\Omega_{x^*} \cap \Omega_{x'}$ as a consequence of q being log-supermodular, along with the fact that $\theta_0 \in \Omega_{x^*} \cap \Omega_{x'}$ and the assumption that h crosses 0 once from below at $\theta = \theta_0$. The second inequality follows from the fact that, for any $\theta \in (\Omega_{x^*} \setminus \Omega_{x'}) \cap (-\infty, \theta_0)$, $h(\theta) \leq 0$, along with the fact that $\Omega_{x^*} \cap (\theta_0, +\infty) = \Omega_{x^*} \cap \Omega_{x'} \cap (\theta_0, \infty)$, with the last property following from noting that the sets Ω_x are ranked in the strong-order sense. The last inequality follows from the observation that, for any $\theta \in \Omega_{x'} \setminus \Omega_{x^*}$, $h(\theta) \geq 0$, which in turn is a consequence of (i) the monotonicity of the sets Ω_x in x , (ii) the assumption that h crosses 0 only once from below at $\theta = \theta_0$, and (iii) the assumption that $\theta_0 \in \Omega_{x^*} \cap \Omega_{x'}$.

Similar arguments imply that, for $x < x^*$, $\Psi(x; \Omega) \leq 0$. The same arguments also imply that, when (a) $\{\theta \in \Omega : h(\theta) \neq 0\}$ has strict positive Lebesgue measure and (b) q is strictly log-supermodular over $\mathbb{R}^2$, then $\Psi(x; \Omega) < 0$ for all $x < x^*$ and $\Psi(x; \Omega) > 0$ for all $x > x^*$. This completes the proof of Property SCB. $\square$

The facts that (a) the continuation game is supermodular, (b) the density $p(x|\theta)$ is log-supermodular, and (c) when agents follow monotone strategies, the regime outcome is monotone in θ imply that, for any $s \in \mathcal{S}$, there exists a unique sequence $(\xi_{(n)}^{\Gamma; s})_{n \in \mathbb{N}}$ such that, for any $n \geq 1$, $\mathbf{a}_{(n)}^{\Gamma}$ is such that $a_{(n), i}^{\Gamma}(x, s) = \mathbf{1}(x > \xi_{(n)}^{\Gamma; s})$ for all i and all $(x, s) \in \mathbb{R} \times \mathcal{S}$, 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 $\xi_{(n)}^{\Gamma; s} \equiv \sup\{x : U^{\Gamma}(x, s | \xi_{(n-1)}^{\Gamma; s}) \leq 0\}$. Indeed, Condition FB together with Property SCB jointly imply that $U^{\Gamma}(x, s | \xi_{(n-1)}^{\Gamma; s}) = \int u(\theta, 1 - P(\xi_{(n-1)}^{\Gamma; s} | \theta)) \Lambda^{\Gamma}(d\theta | x, s)$

crosses zero once from below in x and, therefore, $U^\Gamma(x, s | \xi_{(n-1)}^{\Gamma;s}) > 0$ if, and only if, $x > \xi_{(n)}^{\Gamma;s}$.

Let $T^\Gamma \equiv \cap_{n=1}^\infty T_n^\Gamma$ denote the set of strategy profiles that survive IDISDS under Γ . The most aggressive strategy profile in T^Γ is then given by $a_i^\Gamma(x, s) \equiv \mathbf{1}(x > \xi^{\Gamma;s})$ for all i and all $(x, s) \in \mathbb{R} \times \mathcal{S}$, 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 $\xi^{\Gamma;s} = \sup\{x : U^\Gamma(x, s | x) \leq 0\}$ if $\{x : U^\Gamma(x, s | x) \leq 0\} \neq \emptyset$, and $\xi^{\Gamma;s} \equiv -\infty$ otherwise. This establishes the first part of the lemma. That the profile $\mathbf{a}^\Gamma$ 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 invest for $x_i > \xi^{\Gamma;s}$ and to refrain from investing for $x_i < \xi^{\Gamma;s}$. This completes the proof of the lemma. $\square$

Step 2. Now take any regular policy $\Gamma = (\mathcal{S}, \pi)$ satisfying the perfect-coordination property. Given the result in Theorem 1, without loss of generality, assume that $\Gamma = (\mathcal{S}, \pi)$ is such that $\mathcal{S} = \{0, 1\} \times \hat{\mathcal{S}}$, for some measurable set $\hat{\mathcal{S}}$, and is such that (a) when the policy discloses any signal $s = (\hat{s}, 1)$, all agents invest and default does not happen, whereas (b) when the policy discloses any signal $s = (\hat{s}, 0)$, all agents refrain from investing and default happens.

Equipped with the result in Lemma 1, we show that, starting from $\Gamma = (\mathcal{S}, \pi)$, one can construct a binary policy $\Gamma^* = (\{0, 1\}, \pi^*)$ also satisfying the perfect-coordination property and such that the probability of default under Γ^* is the same as under Γ . The policy $\Gamma^* = (\{0, 1\}, \pi^*)$ is such that, for any θ , $\pi^*(1|\theta) = \int_{\hat{\mathcal{S}}} \pi(d(\hat{s}, 1)|\theta)$. That is, for each θ , the binary policy Γ^* recommends to invest with the same total probability as the original policy Γ discloses signals leading all agents to invest.³⁸

We now show that, under Γ^* , when the policy announces that $s = 1$, the unique rationalizable action for each agent is to invest. To see this, for any $(x, 1)$ that are mutually consistent given Γ^* , let $U^{\Gamma^*}(x, 1|k)$ denote the expected payoff differential for any agent with private signal x , when the policy Γ^* announces $s = 1$, and all other agents follow a cut-off strategy with cut-off k .³⁹ From the law of iterated expectations, we have that

$$U^{\Gamma^*}(x, 1|k) = \int_{\hat{\mathcal{S}}} U^\Gamma(x, (\hat{s}, 1)|k) s^\Gamma(d\hat{s}|x, 1) \quad (15)$$

where $s^\Gamma(\cdot|x, 1)$ is the probability measure over $\hat{\mathcal{S}}$ obtained by conditioning on the event $(x, 1)$, under Γ . For any signal $s = (\hat{s}, 1)$ in the range of π , MARP consistent with Γ is such that $a_i^\Gamma(x, (\hat{s}, 1)) = 1$ all $x \in \mathbb{R}$, and all i , meaning that investing is the unique rationalizable action after Γ announces $s = (\hat{s}, 1)$. Lemma 1 in turn implies that, for all $s = (\hat{s}, 1)$ in the range of π , $\hat{s} \in \hat{\mathcal{S}}$, all $k \in \mathbb{R}$, $U^\Gamma(k, (\hat{s}, 1)|k) > 0$. From (15), we then have that, for all $k \in \mathbb{R}$, $U^{\Gamma^*}(k, 1|k) > 0$. In turn, this implies that, given the new policy Γ^* , when $s = 1$ is disclosed, under MARP consistent with Γ^* , all agents invest, that is, $a_i^{\Gamma^*}(x, 1) = 1$ all x , all $i \in [0, 1]$. It is also easy to see that, when the policy Γ^* discloses the signal $s = 0$, it becomes common certainty among the agents that $\theta \leq \bar{\theta}$. Hence, under MARP consistent

³⁸ $\int_{\hat{\mathcal{S}}} \pi(d(\hat{s}, 1)|\theta)$ represents the total probability that the measure $\pi(\theta)$ assigns to signal $(\hat{s}, 1)$.

³⁹ Recall that $(x, 1)$ are mutually consistent under Γ^* if $p^{\Gamma^*}(x, 1) \equiv \int p(x|\theta)\pi^*(1|\theta) dF(\theta) > 0$.

with Γ^* , after $s = 0$ is disclosed, all agents refrain from investing, irrespective of their private signals. The new policy Γ^* so constructed thus (a) satisfies the perfect-coordination property, and (b) is such that, for any θ , the probability of default under Γ^* is the same as under Γ . $\square$

PROOF OF THEOREM 3*. The conditions in the theorem imply that Theorems 1* and 2* hold. Thus, assume that the policy $\Gamma = (S, \pi)$ (a) is a regular (possibly stochastic) “pass/fail” policy (i.e., $S = \{0, 1\}$, with $\pi(1|\theta) = 1 - \pi(0|\theta)$ denoting the probability that signal $s = 1$ is disclosed when the fundamentals are θ), (b) is such that $\pi(1|\theta) = 0$ for all $\theta \leq \underline{\theta}$ and $\pi(1|\theta) = 1$ for all $\theta > \bar{\theta}$, and (c) satisfies the perfect-coordination property. Theorems 1* and 2* imply that, if Γ does not satisfy these properties, there exists another policy Γ' that satisfies these properties and yields the policymaker a payoff weakly higher than Γ . The proof then follows from applying the arguments below to Γ' instead of Γ .

Suppose that Γ is such that there exists no $\hat{\theta}$ such that $\pi(1|\theta) = 0$ for F -almost all $\theta \leq \hat{\theta}$ and $\pi(1|\theta) = 1$ for F -almost all $\theta > \hat{\theta}$.⁴⁰ We establish the result by showing that there exists a deterministic monotone policy $\Gamma^{\hat{\theta}} = (\{0, 1\}, \pi^{\hat{\theta}})$ satisfying the perfect-coordination property that yields the policymaker a payoff strictly higher than Γ .

Recall that, for the policy Γ to satisfy the perfect-coordination property, it must be that, when the policy discloses the signal $s = 1$, $U^{\Gamma}(x, 1|x) > 0$ for all x such that $(x, 1)$ are mutually consistent, where $U^{\Gamma}(x, 1|x)$ is the expected payoff differential of an agent with signal x who hears that $s = 1$ and who expects all other agents to follow a cut-off strategy with threshold x .

Let $\mathbb{G}$ denote the set of policies $\Gamma' = (S, \pi')$ that, in addition to properties (a) and (b) above, are such that $U^{\Gamma'}(x, 1|x) \geq 0$ for all x such that $(x, 1)$ are mutually consistent.⁴¹ For any $\Gamma \in \mathbb{G}$, let $\mathcal{U}^P[\Gamma]$ denote the policymaker's ex ante expected payoff when, under Γ , agents invest after hearing that $s = 1$ and refrain from investing after hearing that $s = 0$. Denote by $\arg \max_{\tilde{\Gamma} \in \mathbb{G}} \mathcal{U}^P[\tilde{\Gamma}]$ the set of policies that maximize the policymaker's payoff over $\mathbb{G}$.⁴²

Step 1 below shows that any $\Gamma \in \arg \max_{\tilde{\Gamma} \in \mathbb{G}} \mathcal{U}^P[\tilde{\Gamma}]$ is such that $\pi(1|\theta) = 0$ for F -almost all $\theta \leq \theta^*$ and $\pi(1|\theta) = 1$ for F -almost all $\theta > \theta^*$, with θ^* as defined in (4). Step 2 then shows that the policymaker's payoff under the optimal monotone policy $\Gamma^{\theta^*} = (\{0, 1\}, \pi^{\theta^*})$ with cut-off θ^* can be approximated arbitrarily well by a deterministic monotone policy $\Gamma^{\hat{\theta}} = (\{0, 1\}, \pi^{\hat{\theta}}) \in \mathbb{G}$ that satisfies the perfect-coordination property, thus establishing the theorem.

Step 1. Given any policy Γ , let

$$X^{\Gamma} \equiv \{x : (x, 1) \text{ } \Gamma\text{-mutually consistent and } U^{\Gamma}(x, 1|x) = 0\}.$$

⁴⁰If this not the case, then the deterministic monotone policy $\Gamma^{\hat{\theta}} = (\{0, 1\}, \pi^{\hat{\theta}})$ with cut-off $\hat{\theta}$ also satisfies the perfect-coordination property and yields the policymaker the same payoff as Γ , in which case the result trivially holds.

⁴¹As explained in the main text, some policies Γ' in $\mathbb{G}$ need not satisfy the perfect-coordination property, namely those for which there exists x , with $(x, 1)$ mutually consistent, such that $U^{\Gamma'}(x, 1|x) = 0$.

⁴²That $\arg \max_{\tilde{\Gamma} \in \mathbb{G}} \mathcal{U}^P[\tilde{\Gamma}] \neq \emptyset$ follows from the compactness of $\mathbb{G}$ and the upper hemicontinuity of $\mathcal{U}^P$.

Take any policy $\Gamma' \in \mathbb{G}$ for which there exists no $\hat{\theta}$ such that $\pi'(1|\theta) = 0$ for F -almost all $\theta \leq \hat{\theta}$ and $\pi'(1|\theta) = 1$ for F -almost all $\theta > \hat{\theta}$. Clearly, if $X^{\Gamma'} = \emptyset$, there exists another policy $\Gamma'' \in \mathbb{G}$ that yields the policymaker a payoff strictly higher than Γ' .⁴³ Thus, assume that $X^{\Gamma'} \neq \emptyset$, and let $\bar{x} \equiv \sup X^{\Gamma'}$. Claim A below shows that the set $\{\theta \in \Theta(\bar{x}) : \pi'(1|\theta) < 1\}$ has strict positive F -measure. Claim B shows that, given any $\Gamma' \in \mathbb{G}$ for which the posterior beliefs of the marginal agent with signal $\bar{x}$ differ from those obtained by Bayes' rule conditioning on the event that fundamentals are above some threshold $\hat{\theta}$, there exists another policy $\Gamma'' \in \mathbb{G}$ that yields the policymaker a payoff strictly higher than Γ' . Finally, Claim C shows that, under the properties in Condition M*, the only policies $\Gamma' \in \mathbb{G}$ that generate posterior beliefs for the marginal agents with signal $\bar{x}$ equal to those obtained from Bayes' rule by conditioning on the event that fundamentals are above some threshold $\hat{\theta}$ are such that $\pi'(1|\theta) = 0$ for F -almost all $\theta \leq \theta^*$ and $\pi'(1|\theta) = 1$ for F -almost all $\theta > \theta^*$. Jointly, the three claims thus establish the result that any policy $\Gamma \in \arg \max_{\Gamma \in \mathbb{G}} \mathcal{U}^P[\tilde{\Gamma}]$, is such that $\pi(1|\theta) = 0$ for F -almost all $\theta \leq \theta^*$ and $\pi(1|\theta) = 1$ for F -almost all $\theta > \theta^*$.

Given any x , let $\theta_0(x)$ be the fundamental threshold below which the agents' expected payoff differential is negative and above which it is positive, when all agents follow a cut-off strategy with cut-off x . Because Condition FB holds, $\theta_0(x)$ is well-defined.⁴⁴ For any policy $\Gamma = (\{0, 1\}, \pi) \in \mathbb{G}$, let $p^\Gamma(x, 1) \equiv \int_{-\infty}^{+\infty} \pi(1|\theta) p(x|\theta) dF(\theta)$ denote the joint probability density of the exogenous signal x and the endogenous signal $s = 1$.

CLAIM A. *For any $\Gamma' = (\{0, 1\}, \pi') \in \mathbb{G}$ such that $X^{\Gamma'} \neq \emptyset$, $\{\theta \in \Theta(\bar{x}) : \pi'(1|\theta) < 1\}$ has strict positive F -measure.*

PROOF OF CLAIM A. Suppose, by contradiction, that $\pi'(1|\theta) = 1$ for F -almost all $\theta \in \Theta(\bar{x})$. Property (i*) in Condition M* then implies that $\bar{x} > x_{\max}$, where $x_{\max}$ is defined as in (8). In fact, if this was not the case, the monotonicity of $\Theta(\cdot)$ would imply that $\inf \Theta(\bar{x}) \leq \inf \Theta(x_{\max}) < \underline{\theta}$. That $\pi'(1|\theta) = 1$ for F -almost all $\theta \in \Theta(\bar{x})$ would then imply that $\pi'(1|\theta) = 1$ for a set of fundamentals $\theta < \underline{\theta}$ of strict positive F -measure, which is inconsistent with the assumption that $\Gamma' \in \mathbb{G}$. Thus, necessarily, $\bar{x} > x_{\max}$.

Now suppose that $\inf \Theta(\bar{x}) \geq \underline{\theta}$. That $\pi'(1|\theta) = 1$ for F -almost all $\theta \in \Theta(\bar{x})$ means that, from the perspective of an agent with signal $\bar{x}$, the information conveyed by the announcement that $s = 1$ under Γ' is the same as under the monotone deterministic policy $\Gamma^{\underline{\theta}} = (\{0, 1\}, \pi^{\underline{\theta}})$ with cut-off $\hat{\theta} = \underline{\theta}$. As a result, $U^{\Gamma'}(\bar{x}, 1|\bar{x}) = U^{\Gamma^{\underline{\theta}}}(\bar{x}, 1|\bar{x})$. Because $\bar{x} > x_{\max}$, and because, by definition of $x_{\max}$, $U^{\Gamma^{\underline{\theta}}}(x, 1|x) > 0$ for all $x > x_{\max}$, it must be that $U^{\Gamma'}(\bar{x}, 1|\bar{x}) > 0$, which contradicts the assumption that $\bar{x} \in X^{\Gamma'}$. Hence, it must be that $\inf \Theta(\bar{x}) < \underline{\theta}$. As explained above, however, this is inconsistent with the assumption that $\Gamma' \in \mathbb{G}$. $\square$

⁴³In fact, because there exists no such a $\hat{\theta}$, there must exist a set $(\theta', \theta'') \subseteq [\underline{\theta}, \bar{\theta}]$ of F -positive measure over which $\pi'(1|\theta) < 1$. The policy Γ'' can then be obtained from Γ' by increasing $\pi'(1|\theta)$ over such a set. Provided the increase is small, $\Gamma'' \in \mathbb{G}$. Because $U^P(\theta, 1) > U^P(\theta, 0)$ over $[\underline{\theta}, \bar{\theta}]$, the policymaker's payoff under Γ'' is strictly higher than under Γ' .

⁴⁴When the regime outcome is a function of A and θ only, as in the baseline model, $\theta_0(x)$ coincides with the threshold below which default occurs and above which it does not occur when agents follow a cut-off strategy with cut-off x .

Next, for any $\Gamma' = (\{0, 1\}, \pi') \in \mathbb{G}$, let

$$\theta_H \equiv \sup\{\theta \in \Theta : \exists \delta > 0 \text{ s.t. } \pi'(1|\theta') < 1 \text{ for } F\text{-almost all } \theta' \in [\theta - \delta, \theta]\}.$$

The result in Claim A above implies that θ_H is such that $\theta_H > \inf \Theta(\bar{x})$.

CLAIM B. Take any $\Gamma' = (\{0, 1\}, \pi') \in \mathbb{G}$ such that $X^{\Gamma'} \neq \emptyset$. Suppose that

$$\{\theta \in (\underline{\theta}, \theta_H) : \pi'(1|\theta) > 0\} \text{ has strict positive } F\text{-measure.} \quad (16)$$

Then there exists another policy $\Gamma'' \in \mathbb{G}$ that yields the policymaker a payoff strictly higher than Γ' .

Claim B essentially says that, if $\Gamma' \in \mathbb{G}$ is not a deterministic monotone rule, and there exists a $\bar{x}$ such that $U^{\Gamma'}(\bar{x}, 1|\bar{x}) = 0$, then it is improvable.

PROOF OF CLAIM B. The proof below distinguishes two cases.

Case 1: $\underline{\theta} < \theta_0(\bar{x}) \leq \theta_H$. Consider the policy $\Gamma^{\epsilon, \delta} = (\{0, 1\}, \pi^{\epsilon, \delta})$ defined by $\pi^{\epsilon, \delta}(1|\theta) = \pi'(1|\theta)$ for all $\theta \leq \theta_0(\bar{x} + \delta)$, with $\delta > 0$ small so that $\theta_0(\bar{x} + \delta) < \theta_H$, and $\pi^{\epsilon, \delta}(1|\theta) = \min\{\pi'(1|\theta) + \epsilon, 1\}$ for all $\theta > \theta_0(\bar{x} + \delta)$, with $\epsilon > 0$ also small. To see that, when ϵ and δ are small, $\Gamma^{\epsilon, \delta} \in \mathbb{G}$, note that by definition of $\theta_0(\cdot)$, for any x , and any $\theta > \theta_0(x)$, $u(\theta, 1 - P(x|\theta)) > 0$. This property, together with the monotonicity of $\theta_0(\cdot)$, jointly imply that, for any $x \leq \bar{x} + \delta$,

$$\begin{aligned} & \int_{-\infty}^{\infty} u(\theta, 1 - P(x|\theta))(\pi'(1|\theta)\mathbf{1}(\theta \leq \theta_0(\bar{x} + \delta)) \\ & \quad + \min\{\pi'(1|\theta) + \epsilon, 1\}\mathbf{1}(\theta > \theta_0(\bar{x} + \delta)))p(x|\theta) dF(\theta) \\ & \geq \int_{-\infty}^{\infty} u(\theta, 1 - P(x|\theta))\pi'(1|\theta)p(x|\theta) dF(\theta). \end{aligned} \quad (17)$$

To see what justifies the inequality, observe that $u(\theta, 1 - P(\bar{x} + \delta|\theta)) > 0$ for $\theta > \theta_0(\bar{x} + \delta)$, by definition of $\theta_0(\cdot)$. Because, for any θ , $u(\theta, 1 - P(x|\theta))$ is decreasing in x , we then have that, for any $x \leq \bar{x} + \delta$, $u(\theta, 1 - P(x|\theta)) > 0$ for all $\theta > \theta_0(\bar{x} + \delta)$. Because $\Gamma' \in \mathbb{G}$, the right-hand side of (17) is nonnegative.⁴⁵ Hence, for any $x \leq \bar{x} + \delta$ such that $(x, 1)$ are mutually consistent under $\Gamma^{\epsilon, \delta}$, because the left-hand side of (17) is equal to $U^{\Gamma^{\epsilon, \delta}}(x, 1|x)p^{\Gamma^{\epsilon, \delta}}(x, 1)$ and because, for such x , $p^{\Gamma^{\epsilon, \delta}}(x, 1) > 0$, we have that $U^{\Gamma^{\epsilon, \delta}}(x, 1|x) \geq 0$. That $U^{\Gamma^{\epsilon, \delta}}(x, 1|x) \geq 0$ also for all $x > \bar{x} + \delta$ such that $(x, 1)$ are mutually consistent under $\Gamma^{\epsilon, \delta}$ follows from the fact that, by definition of $\bar{x}$, for any $x \geq \bar{x} + \delta$, the function $J(x) \equiv \int_{-\infty}^{+\infty} u(\theta, 1 - P(x|\theta))\pi'(1|\theta)p(x|\theta) dF(\theta)$ is bounded away from 0, along with the fact that, for any $\delta > 0$, the function family $(J^{\epsilon, \delta}(\cdot))_{\epsilon}$ whose elements $J^{\epsilon, \delta}(\cdot)$ are given by $J^{\epsilon, \delta}(x) \equiv \int_{-\infty}^{+\infty} u(\theta, 1 - P(x|\theta))\pi^{\epsilon, \delta}(1|\theta)p(x|\theta) dF(\theta)$ is continuous in ϵ in the sup-norm in a neighborhood of 0.⁴⁶ Because the new policy $\Gamma^{\epsilon, \delta} \in \mathbb{G}$ is such that

⁴⁵Either $(x, 1)$ are not mutually consistent under Γ' , in which case the right-hand side of (17) is zero, or they are mutually consistent, in which case the right-hand side of (17) is equal to $U^{\Gamma'}(x, 1|x)p^{\Gamma'}(x, 1)$, which is nonnegative because $p^{\Gamma'}(x, 1) > 0$ and $U^{\Gamma'}(x, 1|x) \geq 0$.

⁴⁶That is, $\forall k > 0, \exists \Delta > 0$ so that $\forall 0 < \epsilon < \Delta, |J^{\epsilon, \delta}(x) - J(x)| \leq k, \forall x \geq \bar{x} + \delta$.

$\pi^{\epsilon, \delta}(1|\theta) \geq \pi'(1|\theta)$ for all θ , with the inequality strict over a set of fundamentals $\theta \in (\underline{\theta}, \bar{\theta}]$ of F -positive measure, the policymaker's payoff under $\Gamma^{\epsilon, \delta}$ is strictly higher than under Γ' , as claimed.

Case 2: $\theta_H < \theta_0(\bar{x})$. Consider the monotone deterministic policy $\Gamma^\theta = \{\{0, 1\}, \pi^\theta\}$ with cut-off $\hat{\theta} = \underline{\theta}$. Then, for any $x \geq \bar{x}$,

$$\begin{aligned} & \int_{-\infty}^{+\infty} u(\theta, 1 - P(x|\theta)) \pi^\theta(1|\theta) p(x|\theta) dF(\theta) \\ & < \int_{-\infty}^{+\infty} u(\theta, 1 - P(x|\theta)) \pi'(1|\theta) p(x|\theta) dF(\theta), \end{aligned} \quad (18)$$

where the inequality follows from the following facts: (i) $\pi^\theta(1|\theta) = \pi'(1|\theta)$ for F -almost all $\theta \in (-\infty, \underline{\theta}] \cup [\theta_H, +\infty)$ and (ii) $\pi^\theta(1|\theta) = 1 \geq \pi'(1|\theta)$ for F -almost all $\theta \in (\underline{\theta}, \theta_H)$, with the inequality strict over a set of fundamentals in $(\underline{\theta}, \theta_H)$ of strictly positive measure under F , and (iii) $u(\theta, 1 - P(x|\theta)) < 0$ for $\theta \in (\underline{\theta}, \theta_H)$ (by the fact that $\theta_H < \theta_0(\bar{x}) \leq \theta_0(x)$) along with the definition and monotonicity of the function $\theta_0(\cdot)$.

Furthermore, $(\bar{x}, 1)$ are mutually consistent under Γ' , that is, $p^{\Gamma'}(\bar{x}, 1) > 0$. Because $\pi^\theta(1|\theta) \geq \pi'(1|\theta)$ for all θ , $(\bar{x}, 1)$ are mutually consistent also under Γ^θ , that is, $p^{\Gamma^\theta}(\bar{x}, 1) > 0$. Observe that, when $x = \bar{x}$, the left-hand side of (18) is equal to $U^{\Gamma^\theta}(\bar{x}, 1|\bar{x}) p^{\Gamma^\theta}(\bar{x}, 1)$ whereas the right-hand side is equal to $U^{\Gamma'}(\bar{x}, 1|\bar{x}) p^{\Gamma'}(\bar{x}, 1)$. By the definition of $\bar{x}$, $U^{\Gamma'}(\bar{x}, 1|\bar{x}) = 0$, which then implies that $U^{\Gamma^\theta}(\bar{x}, 1|\bar{x}) < 0$. By continuity of $U^{\Gamma^\theta}(x, 1|x)$ in x and the definition of $x_{\max}$ we thus have that $\bar{x} < x_{\max}$. This property in turn permits us to apply Properties (i*) and (ii*) of Condition **M*** below.

Next, let

$$\theta_L \equiv \inf\{\theta \in \Theta : \exists \delta > 0 \text{ s.t. } \pi'(1|\theta) > 0, \text{ with } \pi'(1|\cdot) \text{ continuous over } [\theta, \theta + \delta)\}.$$

By assumption, $\{\theta \in (\underline{\theta}, \theta_H) : \pi'(1|\theta) > 0\}$ has strict positive F -measure. If $\theta_L \geq \theta_H$, then there exists another policy Γ'' for which $\theta_L < \theta_H$ and such that (a) the policymaker's payoff under Γ'' is the same as under Γ' and (b) $U^{\Gamma''}(x, 1|x) = U^{\Gamma'}(x, 1|x)$ for all x . The claim (and ultimately the theorem) then follows from applying the arguments below to Γ'' instead of Γ' . Thus, assume that $\theta_L < \theta_H$. Furthermore, note that $u(\theta_L, 1 - P(\bar{x}|\theta_L)) < 0$.⁴⁷ Also observe that $\inf \Theta(\bar{x}) < \theta_L$. This follows from the fact that, as shown above, $\bar{x} < x_{\max}$, which together with Property (i*) in Condition **M*** and the monotonicity of $\Theta(\cdot)$ in x implies that $\inf \Theta(\bar{x}) < \underline{\theta}$. Because $\theta_L \geq \underline{\theta}$, we thus have that $\inf \Theta(\bar{x}) < \theta_L$.

Recall that $\bar{x}$ is the largest solution to $U^{\Gamma'}(x, 1|x) = 0$. This property, together with the fact that $\Gamma' \in \mathbb{G}$ implies that $U^{\Gamma'}(x, 1|x) > 0$ for all $x > \bar{x}$ such that $(x, 1)$ are mutually consistent under Γ' . Next, observe that, for all $x \geq \bar{x}$, $(x, 1)$ are mutually consistent under Γ' . Because $u(\theta, 1 - P(\bar{x}|\theta)) > 0$ for $\theta > \theta_0(\bar{x})$ and $u(\theta, 1 - P(\bar{x}|\theta)) < 0$ for $\theta < \theta_0(\bar{x})$ (by definition of $\theta_0(\bar{x})$) and because $U^{\Gamma'}(\bar{x}, 1|\bar{x}) = 0$, it must be that $\sup \Theta(\bar{x}) \geq \theta_0(\bar{x})$. Because, by definition of case (2), $\theta_0(\bar{x}) > \theta_H$, this means that $\sup \Theta(\bar{x}) > \theta_H$. The monotonicity of $\Theta(\cdot)$ implies that $\sup \Theta(x) > \theta_H$ for all $x \geq \bar{x}$. Because $\pi'(1|\theta) = 1$ for F -almost all $\theta > \theta_H$, we thus have, for all $x \geq \bar{x}$, $(x, 1)$ are mutually consistent under Γ' (and hence $U^{\Gamma'}(x, 1|x)$ is well-defined for all such x).

⁴⁷This follows from the definition of $\theta_0(\bar{x})$, along with Condition **FB**, and the fact that $\theta_L < \theta_H < \theta_0(\bar{x})$.

The continuity of $U^{\Gamma'}(\cdot, 1|\cdot)$ in x implies that, for *any* $\eta \in (0, x_{\max} - \bar{x})$, the function $U^{\Gamma'}(\cdot, 1|\cdot)$ is bounded away from zero over $[\bar{x} + \eta, x_{\max}]$. That is, there exists $K > 0$ such that $U^{\Gamma'}(x, 1|x) > K$ for all $x \in [\bar{x} + \eta, x_{\max}]$. This property, along with (a) the continuity of $U^{\Gamma'}(\cdot, 1|\cdot)$ in x and (b) the fact that $U^{\Gamma'}(\bar{x}, 1|\bar{x}) = 0$ in turn imply that there exists $\eta \in (0, x_{\max} - \bar{x})$ such that $U^{\Gamma'}(x, 1|x) \geq U^{\Gamma'}(\bar{x} + \eta, 1|\bar{x} + \eta) > 0$ for all $x \in [\bar{x} + \eta, x_{\max}]$.

Now fix $\eta \in (0, x_{\max} - \bar{x})$ such that

$$U^{\Gamma'}(x, 1|x) \geq U^{\Gamma'}(\bar{x} + \eta, 1|\bar{x} + \eta) > 0 \quad \forall x \in [\bar{x} + \eta, x_{\max}]. \quad (19)$$

For any $\epsilon > 0$ small, then let $\delta(\epsilon)$ be implicitly defined by

$$\begin{aligned} & \int_{\theta_L}^{\theta_L + \epsilon} u(\theta, 1 - P(\bar{x} + \eta|\theta)) \pi'(1|\theta) p(\bar{x} + \eta|\theta) dF(\theta) \\ &= \int_{\theta_H - \delta}^{\theta_H} u(\theta, 1 - P(\bar{x} + \eta|\theta)) (1 - \pi'(1|\theta)) p(\bar{x} + \eta|\theta) dF(\theta). \end{aligned} \quad (20)$$

Observe that, for $\epsilon > 0$ small, $\delta(\epsilon)$ is also small, and such that

$$\theta_L + \epsilon < \theta_H - \delta(\epsilon). \quad (21)$$

Also, note that $u(\theta, 1 - P(\bar{x} + \eta|\theta)) < 0$ for all $\theta \in [\theta_L, \theta_H]$. This follows from the fact that $u(\theta, 1 - P(\bar{x} + \eta|\theta)) > 0$ only for $\theta \geq \theta_0(\bar{x} + \eta) > \theta_0(\bar{x}) > \theta_H$.

Consider the policy $\Gamma^{\epsilon, \eta} = \{\{0, 1\}, \pi^{\epsilon, \eta}\}$ defined by the following properties: (a) $\pi^{\epsilon, \eta}(1|\theta) = \pi'(1|\theta)$ for all $\theta \notin \{[\theta_L, \theta_L + \epsilon] \cup [\theta_H - \delta(\epsilon), \theta_H]\}$; (b) $\pi^{\epsilon, \eta}(1|\theta) = 0$ for all $\theta \in [\theta_L, \theta_L + \epsilon]$; and (c) $\pi^{\epsilon, \eta}(1|\theta) = 1$ for all $\theta \in [\theta_H - \delta(\epsilon), \theta_H]$. Because, for any $x \geq \bar{x}$, $(\bar{x}, 1)$ are mutually consistent under Γ' , they are also mutually consistent under $\Gamma^{\epsilon, \eta}$. This follows from the fact that $\pi^{\epsilon, \eta}(1|\theta) = 1$ for F -almost all $\theta > \theta_H$ along with the fact that $\sup \Theta(x) > \theta_H$ for all $x \geq \bar{x}$, as shown above. Hence, $U^{\Gamma^{\epsilon, \eta}}(x, 1|x)$ is well-defined for all $x \geq \bar{x}$. Also, observe that $p^{\Gamma^{\epsilon, \eta}}(\bar{x} + \eta, 1)$ need not coincide with $p^{\Gamma'}(\bar{x} + \eta, 1)$. However, Condition (20) implies that

$$U^{\Gamma^{\epsilon, \eta}}(\bar{x} + \eta, 1|\bar{x} + \eta) \stackrel{\text{sgn}}{=} U^{\Gamma'}(\bar{x} + \eta, 1|\bar{x} + \eta) > 0.$$

We now show, for any $\eta \in (0, x_{\max} - \bar{x})$ satisfying Condition (19), $\epsilon > 0$ satisfying Condition (21), and x such that $(x, 1)$ are mutually consistent under $\Gamma^{\epsilon, \eta}$, $U^{\Gamma^{\epsilon, \eta}}(x, 1|x) \geq 0$. Recall that, by the definition of $x_{\max}$, for all $x > x_{\max}$, $U^{\Gamma^{\theta}}(x, 1|x)$ is well-defined and strictly positive, implying that

$$\int_{-\infty}^{+\infty} u(\theta, 1 - P(x|\theta)) \pi^{\theta}(1|\theta) p(x|\theta) dF(\theta) > 0.$$

Also recall that, for any x , the payoff differential $u(\theta, 1 - P(x|\theta))$ is negative for $\theta < \theta_0(x)$ and positive for $\theta > \theta_0(x)$, and that, for any $x > x_{\max}$, $\theta_0(x) > \theta_0(x_{\max}) > \theta_0(\bar{x}) > \theta_H$.

Because the policy $\Gamma^{\epsilon, \eta}$ is such that $\pi^{\epsilon, \eta}(1|\theta) = \pi^{\theta}(1|\theta)$ for all $\theta > \theta_H$ and $\pi^{\epsilon, \eta}(1|\theta) \leq \pi^{\theta}(1|\theta)$ for all $\theta < \theta_H$, with the inequality strict over a set of strictly positive measure

under F , we have that

$$\begin{aligned} & \int_{-\infty}^{+\infty} u(\theta, 1 - P(x|\theta)) \pi^{\epsilon, \eta}(1|\theta) p(x|\theta) dF(\theta) \\ & > \int_{-\infty}^{+\infty} u(\theta, 1 - P(x|\theta)) \pi^{\bar{\theta}}(1|\theta) p(x|\theta) dF(\theta). \end{aligned}$$

Because the right-hand side is strictly positive, for any $x > x_{\max}$, $U^{\Gamma^{\epsilon, \eta}}(x, 1|x) > 0$.

Next, for any $\theta \in [\underline{\theta}, \bar{\theta}]$, let $x^*(\theta)$ be the signal threshold such that, when all agents invest for $x > x^*(\theta)$ and refrain from investing for $x < x^*(\theta)$, the expected payoff differential $u(\bar{\theta}, 1 - P(x^*(\theta)|\bar{\theta}))$ is positive if and only if the fundamentals $\bar{\theta}$ are above θ . Observe that, for any $\theta \in [\underline{\theta}, \bar{\theta}]$, the existence of such a threshold follows from Condition FB, and that $x^*(\underline{\theta}) = -\infty$ and $x^*(\bar{\theta}) = +\infty$. Clearly, for any $\eta \in (0, x_{\max} - \bar{x})$ satisfying Condition (19), $\epsilon > 0$ satisfying Condition (21), and $x \leq x^*(\theta_L + \epsilon)$,

$$\int_{\theta_L}^{+\infty} u(\theta, 1 - P(x|\theta)) p(x|\theta) \pi^{\epsilon, \eta}(1|\theta) dF(\theta) > 0.$$

This is because for any $x \leq x^*(\theta_L + \epsilon)$, $\theta_0(x) \leq \theta_L + \epsilon$. The result then follows from the fact that, for any $x \leq x^*(\theta_L + \epsilon)$, $\pi^{\epsilon, \eta}(1|\theta) = 0$ for all $\theta \leq \theta_0(x)$. Hence, for any $x \leq x^*(\theta_L + \epsilon)$ such that $(x, 1)$ are mutually consistent under $\Gamma^{\epsilon, \eta}$, $U^{\Gamma^{\epsilon, \eta}}(x, 1|x) \geq 0$.

Next, observe that for any $x \in (x^*(\theta_L + \epsilon), x^*(\theta_H - \delta(\epsilon))]$,

$$\begin{aligned} & \int_{\theta_L}^{+\infty} u(\theta, 1 - P(x|\theta)) p(x|\theta) \pi^{\epsilon, \eta}(1|\theta) dF(\theta) \\ &= \int_{\theta_L}^{+\infty} u(\theta, 1 - P(x|\theta)) p(x|\theta) \pi'(1|\theta) dF(\theta) \\ & \quad - \int_{\theta_L}^{\theta_L + \epsilon} u(\theta, 1 - P(x|\theta)) p(x|\theta) \pi'(1|\theta) dF(\theta) \\ & \quad + \int_{\theta_H - \delta(\epsilon)}^{\theta_H} u(\theta, 1 - P(x|\theta)) p(x|\theta) (1 - \pi'(1|\theta)) dF(\theta) \geq 0. \end{aligned}$$

To understand the inequality, first observe that the first integral on the right-hand side of the equality is nonnegative (if it was strictly negative then $(x, 1)$ would be mutually consistent under Γ' and $U^{\Gamma'}(x, 1|x) < 0$, which is inconsistent with the fact that $\Gamma' \in \mathbb{G}$). Second, observe that the integrand function in the second integral on the right-hand side of the equality is nonpositive (this follows from the fact that, when $x > x^*(\theta_L + \epsilon)$, $u(\theta, 1 - P(x|\theta)) \leq 0$ for all $\theta \leq \theta_L + \epsilon$). Finally, the integrand function in the third integral on the right-hand side of the equality is nonnegative (this follows from the fact that when $x < x^*(\theta_H - \delta(\epsilon))$, $u(\theta, 1 - P(x|\theta)) \geq 0$ for all $\theta > \theta_H - \delta(\epsilon)$). Hence, for any such x , if $(x, 1)$ are mutually consistent under $\Gamma^{\epsilon, \eta}$, it must be that $U^{\Gamma^{\epsilon, \eta}}(x, 1|x) \geq 0$.

Next, consider $x \in (x^*(\theta_H - \delta(\epsilon)), x^*(\theta_H))$. For any x , let

$$\Delta S(x) \equiv \int_{\theta_L}^{+\infty} u(\theta, 1 - P(x|\theta)) p(x|\theta) (\pi^{\epsilon, \eta}(1|\theta) - \pi'(1|\theta)) dF(\theta),$$

and, for any (x, θ) , let $q(\theta, x) \equiv |u(\theta, 1 - P(x|\theta))|p(x|\theta)$. Note that, for any

$$x \in (x^*(\theta_H - \delta(\epsilon)), x^*(\theta_H)),$$

$\theta_0(x) \in (\theta_H - \delta(\epsilon), \theta_H)$, and

$$\begin{aligned} \Delta S(x) &= \int_{\theta_L}^{\theta_H - \delta(\epsilon)} -u(\theta, 1 - P(x|\theta))p(x|\theta)(\pi'(1|\theta) - \pi^{\epsilon, \eta}(1|\theta)) dF(\theta) \\ &\quad + \int_{\theta_H - \delta(\epsilon)}^{\theta_0(x)} -u(\theta, 1 - P(x|\theta))p(x|\theta)(\pi'(1|\theta) - \pi^{\epsilon, \eta}(1|\theta)) dF(\theta) \\ &\quad + \int_{\theta_0(x)}^{\theta_H} -u(\theta, 1 - P(x|\theta))p(x|\theta)(\pi'(1|\theta) - \pi^{\epsilon, \eta}(1|\theta)) dF(\theta) \\ &\geq \int_{\theta_L}^{\theta_H - \delta(\epsilon)} \frac{q(\theta, x)}{q(\theta, \bar{x} + \eta)} q(\theta, \bar{x} + \eta)(\pi'(1|\theta) - \pi^{\epsilon, \eta}(1|\theta)) dF(\theta) \\ &\quad + \int_{\theta_H - \delta(\epsilon)}^{\theta_0(x)} \frac{q(\theta, x)}{q(\theta, \bar{x} + \eta)} q(\theta, \bar{x} + \eta)(\pi'(1|\theta) - \pi^{\epsilon, \eta}(1|\theta)) dF(\theta) \\ &\quad + \frac{q(\theta_H - \delta(\epsilon), x)}{q(\theta_H - \delta(\epsilon), \bar{x} + \eta)} \int_{\theta_0(x)}^{\theta_H} q(\theta, \bar{x} + \eta)(\pi'(1|\theta) - \pi^{\epsilon, \eta}(1|\theta)) dF(\theta) \\ &\geq \frac{q(\theta_H - \delta(\epsilon), x)}{q(\theta_H - \delta(\epsilon), \bar{x} + \eta)} \int_{\theta_L}^{\theta_H} q(\theta, \bar{x} + \eta)(\pi'(1|\theta) - \pi^{\epsilon, \eta}(1|\theta)) dF(\theta) \\ &= \frac{q(\theta_H - \delta(\epsilon), x)}{q(\theta_H - \delta(\epsilon), \bar{x} + \eta)} \Delta S(\bar{x} + \eta) = 0. \end{aligned}$$

The first equality follows from the definition of the $\Delta S(x)$ function. The first inequality follows from the fact that (i) for any $\theta \leq \theta_0(x)$, $u(\theta, 1 - P(x|\theta)) < 0$, whereas for any $\theta > \theta_0(x)$, $u(\theta, 1 - P(x|\theta)) > 0$, and (ii) for $\theta \in [\theta_0(x), \theta_H]$, $\pi'(1|\theta) \leq \pi^{\epsilon, \eta}(1|\theta)$. Together, these properties imply that

$$\begin{aligned} &\int_{\theta_0(x)}^{\theta_H} -u(\theta, 1 - P(x|\theta))p(x|\theta)(\pi'(1|\theta) - \pi^{\epsilon, \eta}(1|\theta)) dF(\theta) \\ &\geq 0 \geq \frac{q(\theta_H - \delta(\epsilon), x)}{q(\theta_H - \delta(\epsilon), \bar{x} + \eta)} \int_{\theta_0(x)}^{\theta_H} q(\theta, \bar{x} + \eta)(\pi'(1|\theta) - \pi^{\epsilon, \eta}(1|\theta)) dF(\theta). \end{aligned}$$

The second inequality follows from the fact that $\pi'(1|\theta) - \pi^{\epsilon, \eta}(1|\theta)$ turns from positive to negative at $\theta = \theta_H - \delta(\epsilon) \leq \theta_0(x)$, along with the fact that, for any $\theta \in [\theta_L, \theta_0(x)]$, the function $q(\theta, x)/q(\theta, \bar{x} + \eta)$ is nonincreasing in θ as implied by the log-supermodularity of $|u(\theta, 1 - P(x|\theta))|p(x|\theta)$ over $\{(\theta, x) \in [0, 1] \times \mathbb{R} : u(\theta, 1 - P(x|\theta)) \leq 0\}$, which in turn follows from Property (iii*) of Condition **M*** and the assumption that $p(x|\theta)$ is log-supermodular. The second equality follows from the fact that $\theta_0(\bar{x} + \eta) > \theta_0(\bar{x}) > \theta_H$, which implies that $u(\theta, 1 - P(\bar{x} + \eta|\theta)) \leq 0$ for all $\theta \leq \theta_H$. Finally, the last equality follows from the fact that, by construction of the policy $\Gamma^{\epsilon, \eta}$, $\Delta S(\bar{x} + \eta) = 0$. Hence, for any

$x \in (x^*(\theta_H - \delta(\epsilon)), x^*(\theta_H))$, $\Delta S(x) \geq 0$, which implies that, for any x in this range such that $(x, 1)$ are mutually consistent under $\Gamma^{\epsilon, \eta}$, $U^{\Gamma^{\epsilon, \eta}}(x, 1|x) \geq 0$.

Similar arguments imply that, for any $x \in [x^*(\theta_H), \bar{x} + \eta]$,

$$\begin{aligned} \Delta S(x) &= \int_{\theta_L}^{\theta_H} -u(\theta, 1 - P(x|\theta)) p(x|\theta) (\pi'(1|\theta) - \pi^{\epsilon, \eta}(1|\theta)) dF(\theta) \\ &= \int_{\theta_L}^{\theta_H} \frac{q(\theta, x)}{q(\theta, \bar{x} + \eta)} q(\theta, \bar{x} + \eta) (\pi'(1|\theta) - \pi^{\epsilon, \eta}(1|\theta)) dF(\theta) \\ &\geq \frac{q(\theta_H - \delta(\epsilon), x)}{q(\theta_H - \delta(\epsilon), \bar{x} + \eta)} \Delta S(\bar{x} + \eta) = 0, \end{aligned}$$

implying that, for such x , too, if $(x, 1)$ are mutually consistent under $\Gamma^{\epsilon, \eta}$, then $U^{\Gamma^{\epsilon, \eta}}(x, 1|x) \geq 0$ (this result also follows from Property (iii*) of Condition **M*** along with the log-supermodularity of $p(x|\theta)$).

Thus far, we have established that, for any $x \in (-\infty, \bar{x} + \eta) \cup (x_{\max}, +\infty)$ such that $(x, 1)$ are mutually consistent under $\Gamma^{\epsilon, \eta}$, $U^{\Gamma^{\epsilon, \eta}}(x, 1|x) \geq 0$. Below we show that there exists a $\bar{\epsilon} > 0$ such that, for any $\eta \in (0, x_{\max} - \bar{x})$ satisfying Condition (19), and $\epsilon > 0$ satisfying Condition (21), with $\epsilon \in [0, \bar{\epsilon}]$, the same is true also for any $x \in [\bar{x} + \eta, x_{\max}]$ such that $(x, 1)$ are mutually consistent under $\Gamma^{\epsilon, \eta}$. For any x , let

$$S^{\Gamma^{\epsilon, \eta}}(x) \equiv \int_{\theta_L}^{+\infty} u(\theta, 1 - P(x|\theta)) p(x|\theta) \pi^{\epsilon, \eta}(1|\theta) dF(\theta).$$

Note that, for any η , the function family $(S^{\Gamma^{\epsilon, \eta}}(\cdot))_{\epsilon}$ is continuous in ϵ in the sup-norm, in a neighborhood of 0. That is, for any $z > 0$, there exists $\kappa > 0$ such that, for any $0 < \epsilon < \kappa$, and all x , $|S^{\Gamma^{\epsilon, \eta}}(x) - S^{\Gamma^{0, \eta}}(x)| \leq z$, where $\Gamma^{0, \eta} = \Gamma'$.⁴⁸ By Condition (19), $U^{\Gamma'}(x, 1|x)$ is bounded away from zero over $[\bar{x} + \eta, x_{\max}]$. Hence, there exists $\bar{\epsilon} > 0$ small such that, for any $\eta \in (0, x_{\max} - \bar{x})$ satisfying Condition (19), and $\epsilon \in [0, \bar{\epsilon}]$, Condition (21) holds and, for any $x \in [\bar{x} + \eta, x_{\max}]$, $U^{\Gamma^{\epsilon, \eta}}(x, 1|x) > 0$.⁴⁹

Together, the results above thus imply that, for any $\eta \in (0, x_{\max} - \bar{x})$ satisfying Condition (19), and $\epsilon \in [0, \bar{\epsilon}]$, the policy $\Gamma^{\epsilon, \eta} \in \mathbb{G}$.

We now show that, when Property (ii*) in Condition **M*** holds, for any $\eta \in (0, x_{\max} - \bar{x})$ satisfying Condition (19), and $\epsilon \in [0, \bar{\epsilon}]$, the new policy $\Gamma^{\epsilon, \eta}$ yields the policymaker an expected payoff strictly higher than Γ' . To see this, observe that the policymaker's payoff under any such policy is equal to

$$\begin{aligned} \mathcal{U}^P[\Gamma^{\epsilon, \eta}] &= \int_{-\infty}^{\theta_L + \epsilon} U^P(\theta, 0) dF(\theta) + \int_{\theta_H - \delta(\epsilon)}^{\theta_H} U^P(\theta, 1) dF(\theta) \\ &\quad + \int_{(\theta_L + \epsilon, \theta_H - \delta(\epsilon)) \cup (\theta_H, +\infty)} (\pi'(1|\theta) U^P(\theta, 1) + (1 - \pi'(1|\theta)) U^P(\theta, 0)) dF(\theta). \end{aligned}$$

⁴⁸This follows from the fact that $|u(\theta, A)|$ and $p(x|\theta)$ are both bounded.

⁴⁹Recall that, as established above, for any $x \geq \bar{x} + \eta$, $(x, 1)$ are mutually consistent under $\Gamma^{\epsilon, \eta}$.

Differentiating $\mathcal{U}^P[\Gamma^{\epsilon, \eta}]$ with respect to ϵ , and using the implicit function theorem to obtain the derivative of $\delta(\epsilon)$, we have that

$$\begin{aligned} \frac{d\mathcal{U}^P[\Gamma^{\epsilon, \eta}]}{d\epsilon} &= f(\theta_H - \delta)(1 - \pi'(1|\theta_H - \delta))[U^P(\theta_H - \delta, 1) - U^P(\theta_H - \delta, 0)] \times \delta'(\epsilon) \\ &\quad - f(\theta_L + \epsilon)\pi'(1|\theta_L + \epsilon)[U^P(\theta_L + \epsilon, 1) - U^P(\theta_L + \epsilon, 0)] \\ &= f(\theta_L + \epsilon)\pi'(1|\theta_L + \epsilon)[U^P(\theta_H - \delta, 1) - U^P(\theta_H - \delta, 0)] \\ &\quad \times \frac{p(\bar{x} + \eta|\theta_L + \epsilon)u(\theta_L + \epsilon, 1 - P(\bar{x} + \eta|\theta_L + \epsilon))}{p(\bar{x} + \eta|\theta_H - \delta)u(\theta_H - \delta, 1 - P(\bar{x} + \eta|\theta_H - \delta))} \\ &\quad - f(\theta_L + \epsilon)\pi'(1|\theta_L + \epsilon)[U^P(\theta_L + \epsilon, 1) - U^P(\theta_L + \epsilon, 0)]. \end{aligned}$$

Property (ii*) in Condition **M***, together with the fact that $\bar{x} \leq x_{\max}$, guarantee that, for any $\eta \in (0, x_{\max} - \bar{x})$ satisfying Condition (19), and $\epsilon \in (0, \bar{\epsilon}]$, $d\mathcal{U}^P[\Gamma^{\epsilon, \eta}]/d\epsilon > 0$. We conclude that the policy $\Gamma^{\bar{\epsilon}, \eta} \in \mathbb{G}$ yields the policymaker a payoff strictly higher than Γ' . This completes the proof of Claim S1-B. $\square$

CLAIM C. Take any $\Gamma' = (\{0, 1\}, \pi') \in \mathbb{G}$ such that $X^{\Gamma'} \neq \emptyset$ and

$$\{\theta \in (\underline{\theta}, \theta_H) : \pi'(1|\theta) > 0\} \quad \text{has zero } F\text{-measure.} \quad (22)$$

Then $\pi'(1|\theta) = 0$ for F -almost all $\theta \leq \theta^*$ and $\pi'(1|\theta) = 1$ for F -almost all $\theta > \theta^*$.

Claim C says that, if $\Gamma' \in \mathbb{G}$ is a deterministic monotone rule, and there exists a $\bar{x}$ such that $U^{\Gamma'}(\bar{x}, 1|\bar{x}) = 0$, then Γ' differs from the optimal monotone rule Γ^{θ^*} over at most a set of fundamentals of zero F -measure.

PROOF OF CLAIM C. Let $\Gamma^{\theta_H} = (\{0, 1\}, \pi^{\theta_H})$ be the deterministic monotone policy with cut-off θ_H . Clearly, any x such that $(x, 1)$ are mutually consistent under Γ' is such that $(x, 1)$ are also mutually consistent under Γ^{θ_H} . Furthermore, for any such x , $U^{\Gamma'}(x, 1|x) = U^{\Gamma^{\theta_H}}(x, 1|x)$ (both properties follow because the two policies differ only over sets of zero F -measure).

Suppose that $\theta_H > \theta^*$. Below we establish that, in this case, any x such that $(x, 1)$ are mutually consistent under Γ^{θ_H} is such that $U^{\Gamma^{\theta_H}}(x, 1|x) > 0$. Clearly, for any x such that (a) $\theta_0(x) \leq \theta_H$, and (b) $(x, 1)$ are mutually consistent under Γ^{θ_H} , $U^{\Gamma^{\theta_H}}(x, 1|x) > 0$. Thus, consider any x such that $\theta_0(x) > \theta_H$, and (b) $(x, 1)$ are mutually consistent under Γ^{θ_H} . Note first that, for any such x , $(x, 1)$ are mutually consistent also under $\Gamma^{\theta^*} = (\{0, 1\}, \pi^{\theta^*})$. This is because $\pi^{\theta^*}(1|\theta) \geq \pi^{\theta_H}(1|\theta)$ for all θ . Furthermore, for any such x ,

$$\int_{\theta_H}^{+\infty} u(\theta, 1 - P(x|\theta))p(x|\theta) dF(\theta) \geq \int_{\theta^*}^{+\infty} u(\theta, 1 - P(x|\theta))p(x|\theta) dF(\theta). \quad (23)$$

This follows from the fact that $u(\theta, 1 - P(x|\theta)) < 0$ for all $\theta \in [\theta^*, \theta_H]$. Hence, for any such x , because $U^{\Gamma^{\theta^*}}(x, 1|x) \geq 0$, $U^{\Gamma^{\theta_H}}(x, 1|x) \geq 0$. Now take $x = \bar{x}$ and recall that, by definition, $U^{\Gamma'}(\bar{x}, 1|\bar{x}) = 0$. Because $U^{\Gamma'}(\bar{x}, 1|\bar{x}) = U^{\Gamma^{\theta_H}}(\bar{x}, 1|\bar{x})$, this means that

$U^{\Gamma^{\theta_H}}(\bar{x}, 1|\bar{x}) = 0$. Property (i*) of Condition **M*** then implies that $\inf \Theta(\bar{x}) < \underline{\theta}$. Hence, for $x = \bar{x}$, the inequality in (23) is strict, which in turn implies that $U^{\Gamma^{\theta^*}}(\bar{x}, 1|\bar{x}) < 0$, contradicting the assumption that $\Gamma^{\theta^*} \in \mathbb{G}$. Therefore, it must be that $\theta_H \leq \theta^*$. However, by definition of θ^* , if $\theta_H < \theta^*$, there exists an x such that (a) $U^{\Gamma^{\theta_H}}(x, 1|x) < 0$, and (b) $(x, 1)$ are mutually consistent under Γ^{θ_H} . Because, for all such x , $U^{\Gamma'}(x, 1|x)$ is well-defined (i.e., $(x, 1)$ are mutually consistent also under Γ') and $U^{\Gamma^{\theta_H}}(x, 1|x) = U^{\Gamma'}(x, 1|x)$, we thus have that $U^{\Gamma'}(x, 1|x) < 0$, which contradicts the assumption that $\Gamma' \in \mathbb{G}$. Hence, $\theta_H = \theta^*$. This completes the proof of Claim C. $\square$

Step 2. Step 1 implies that $\arg \max_{\tilde{\Gamma} \in \mathbb{G}} \mathcal{U}^P[\tilde{\Gamma}] \neq \emptyset$ and that any $\Gamma^* = (\{0, 1\}, \pi)$ with $\Gamma^* \in \arg \max_{\tilde{\Gamma} \in \mathbb{G}} \mathcal{U}^P[\tilde{\Gamma}]$ is such that $\pi(1|\theta) = 0$ for F -almost all $\theta \leq \theta^*$ and $\pi(1|\theta) = 1$ for F -almost all $\theta > \theta^*$. The result in the theorem then follows from observing that, given any $\Gamma^* \in \arg \max_{\tilde{\Gamma} \in \mathbb{G}} \mathcal{U}^P[\tilde{\Gamma}]$, there exists a nearby deterministic monotone policy $\Gamma^{\hat{\theta}} \in \mathbb{G}$ with cut-off $\hat{\theta} = \theta^* + \tilde{\varepsilon}$, for $\tilde{\varepsilon} > 0$ small, such that $\Gamma^{\hat{\theta}}$ satisfies the perfect-coordination property (i.e., $U^{\Gamma^{\hat{\theta}}}(x, 1|x) > 0$ all x such that $(x, 1)$ are mutually consistent under $\Gamma^{\hat{\theta}}$).⁵⁰ The continuity of $\mathcal{U}^P[\Gamma^{\hat{\theta}}]$ in $\hat{\theta}$ then implies that, for $\tilde{\varepsilon} > 0$ small, $\mathcal{U}^P[\Gamma^{\hat{\theta}}] > \mathcal{U}^P[\Gamma^*]$, thus establishing the result in the theorem. $\square$

REFERENCES

- Alonso, Ricardo and Odilon Camara (2016a), “Persuading voters.” *American Economic Review*, 106, 3590–3605. [0767]
- Alonso, Ricardo and Konstantinos E. Zachariadis (2023), “Persuading large investors.” Available at SSRN 3781499. [0767]
- Alvarez, Fernando and Gadi Barlevy (2021), “Mandatory disclosure and financial contagion.” *Journal of Economic Theory*, 194, 105237. [0768]
- Angeletos, George-Marios, Christian Hellwig, and Alessandro Pavan (2006), “Signaling in a global game: Coordination and policy traps.” *Journal of Political Economy*, 114, 452–484. [0767]
- Angeletos, George-Marios, Christian Hellwig, and Alessandro Pavan (2007), “Dynamic global games of regime change: Learning, multiplicity, and the timing of attacks.” *Econometrica*, 75, 711–756. [0767, 0793]
- Angeletos, George-Marios and Alessandro Pavan (2013), “Selection-free predictions in global games with endogenous information and multiple equilibria.” *Theoretical Economics*, 8, 883–938. [0767]
- Angeletos, George-Marios and Iván Werning (2006), “Crises and prices: Information aggregation, multiplicity, and volatility.” *American Economic Review*, 96, 1720–1736. [0767]
- Arieli, Itai and Yakov Babichenko (2019), “Private Bayesian persuasion.” *Journal of Economic Theory*, 182, 185–217. [0767]

⁵⁰The arguments are the same as those used in the proof of Claim C for the case $\theta_H > \theta^*$.

- Bardhi, Arjada and Yingni Guo (2018), “Modes of persuasion toward unanimous consent.” *Theoretical Economics*, 13, 1111–1149. [0767]
- Basak, Deepal and Zhen Zhou (2020), “Diffusing coordination risk.” *American Economic Review*, 110, 271–297. [0767]
- Basak, Deepal and Zhen Zhou (2022), “Panics and early warnings.” Manuscript, PBCSF-NIFR. [0793]
- Bergemann, Dirk and Stephen Morris (2019), “Information design: A unified perspective.” *Journal of Economic Literature*, 57, 44–95. [0767]
- Bouvard, Matthieu, Pierre Chaigneau, and Adolfo de Motta (2015), “Transparency in the financial system: Rollover risk and crises.” *The Journal of Finance*, 70, 1805–1837. [0767]
- Calzolari, Giacomo and Alessandro Pavan (2006a), “On the optimality of privacy in sequential contracting.” *Journal of Economic theory*, 130, 168–204. [0793]
- Calzolari, Giacomo and Alessandro Pavan (2006b), “Monopoly with resale.” *The RAND Journal of Economics*, 37, 362–375. [0793]
- Che, Yeon-Koo and Johannes Hörner (2018), “Recommender systems as mechanisms for social learning.” *Quarterly Journal of Economics*, 133, 871–925. [0767]
- Cong, Lin William, Steven R. Grenadier, and Yunzhi Hu (2020), “Dynamic interventions and informational linkages.” *Journal of Financial Economics*, 135, 1–15. [0767]
- Corona, Carlos, Lin Nan, and Zhang Gaoqing (2017), “The coordination role of stress test disclosure in bank risk taking.” Manuscript, Carnegie Mellon University. [0768]
- Denti, Tommaso (2023), “Unrestricted information acquisition.” *Theoretical Economics*, 18, 1101–1140. [0767]
- Diamond, Douglas W. and Philip H. Dybvig (1983), “Bank runs, deposit insurance, and liquidity.” *Journal of Political Economy*, 91. [0768]
- Doval, Laura and Jeffrey C. Ely (2020), “Sequential information design.” *Econometrica*, 88, 2575–2608. [0767]
- Dworczak, Piotr (2020), “Mechanism design with aftermarkets: Cutoff mechanisms.” *Econometrica*, 88, 2629–2661. [0793]
- Dworczak, Piotr and Alessandro Pavan (2022), “Preparing for the worst but hoping for the best: Robust (Bayesian) persuasion.” *Econometrica*, 90, 2017–2051. [0793]
- Edmond, Chris (2013), “Information manipulation, coordination, and regime change.” *Review of Economic Studies*, 80, 1422–1458. [0767]
- Ely, Jeffrey C. (2017) *Beeps*. *The American Economic Review*, 107, 31–53. [0793]
- Farhi, Emmanuel and Jean Tirole (2012), “Collective moral hazard, maturity mismatch, and systemic bailouts.” *American Economic Review*, 102, 60–93. [0768]

Faria-e Castro, Miguel, Joseba Martinez, and Thomas Philippon (2016), “Runs versus lemons: Information disclosure and fiscal capacity.” *The Review of Economic Studies*, 060. [0767, 0768]

Flannery, Mark, Beverl Hirtle, and Anna Kovner (2017), “Evaluating the information in the federal reserve stress tests.” *Journal of Financial Intermediation*, 29, 1–18. [0768]

Frankel, David M. (2017), “Efficient ex-ante stabilization of firms.” *Journal of Economic Theory*, 170, 112–144. [0766]

Galperti, Simone and Jacopo Perego (2023), “Games with information constraints: Seeds and spillovers.” Manuscript, Columbia GSB. [0767]

Galvão, Raphael and Felipe Shalders (2022), “Rules versus discretion in central bank communication.” *International Journal of Economic Theory*. [0766, 0767]

Garcia, Filomena and Ettore Panetti (2017), “A theory of government bailouts in a heterogeneous banking union.” Manuscript, Indiana University. [0768]

Gick, Wolfgang and Thilo Pausch (2012), “Persuasion by stress testing: Optimal disclosure of supervisory information in the banking sector.” Manuscript, Bundesbank. [0767]

Gitmez, A. Arda and Pooya Molavi (2023), “Informational Autocrats, Diverse Societies” Manuscript, Northwestern University. [0767]

Goldstein, Itay and Chong Huang (2016), “Bayesian persuasion in coordination games.” *The American Economic Review*, 106, 592–596. [0766, 0767, 0778, 0780]

Goldstein, Itay and Yaron Leitner (2018), “Stress tests and information disclosure.” *Journal of Economic Theory*, 177, 34–69. [0767, 0792]

Goldstein, Itay and Haresh Sapra (2014), “Should banks’ stress test results be disclosed? An analysis of the costs and benefits.” *Foundations and Trends (R) in Finance*, 8, 1–54. [0767]

Guo, Yingni and Eran Shmaya (2019), “The interval structure of optimal disclosure.” *Econometrica*, 87, 653–675. [0791]

Halac, Marina, Ilan Kremer, and Eyal Winter (2020), “Raising capital from heterogeneous investors.” *American Economic Review*, 110, 889–921. [0766]

Halac, Marina, Elliot Lipnowski, and Daniel Rappoport (2021), “Rank uncertainty in organizations.” *American Economic Review*, 111, 757–786. [0766]

Halac, Marina, Elliot Lipnowski, and Daniel Rappoport (2022), “Addressing strategic uncertainty with incentives and information.” *AEA Papers and Proceedings*, 112, 431–437. [0766]

Heese, Carl and Stephan Laueremann (2021), “Persuasion and information aggregation in elections.” Manuscript, University of Bonn. [0767]

Henry, Jerome and Christoffer Kok (2013), “A macro stress testing framework for assessing systemic risks in the banking sector.” Manuscript, European Central Bank. [0764]

Homar, Timotej, Heinrich Kick, and Carmelo Salleo (2016), “Making sense of the EU wide stress test: A comparison with the srisk approach.” Manuscript, European Central Bank. [0764]

Inostroza, Nicolas (2023), “Persuading multiple audiences: Strategic complementarities and (robust) regulatory disclosures.” Manuscript, University of Toronto. [0768]

Inostroza, Nicolas and Alessandro Pavan (2024a), “Adversarial coordination and public information design: Additional material.” Manuscript, Northwestern University. [0766, 0788]

Inostroza, Nicolas and Alessandro Pavan (2024b), “Bank funding and optimal stress tests.” Manuscript, Northwestern University. [0768]

Kamenica, Emir (2019), “Bayesian persuasion and information design.” *Annual Review of Economics*. (forthcoming). [0767]

Kyriazis, Panagiotis and Edmund Lou (2024), “(in)effectiveness of information manipulation in regime-change games.” Manuscript, Northwestern University. [0767]

Laclau, Marie and Ludovic Renou (2017), “Public persuasion.” Manuscript, HEC Paris. [0767]

Leitner, Yaron and Basil Williams (2023), “Model secrecy and stress tests.” *The Journal of Finance*, 78, 1055–1095. [0770]

Lerner, Josh and Jean Tirole (2006), “A model of forum shopping.” *American economic review*, 96, 1091–1113. [0764]

Li, Fei, Yangbo Song, and Mofei Zhao (2023), “Global manipulation by local obfuscation.” *Journal of Economic Theory*, 207, 105575. [0766]

Mathevet, Laurent, Jacopo Perego, and Ina Taneva (2020), “On information design in games.” *Journal of Political Economy*, 128, 1370–1404. [0767]

Mensch, Jeffrey (2021), “Monotone persuasion.” *Games and Economic Behavior*, 130, 521–542. [0792]

Morgan, Donald, Stravos Persitani, and Vanessa Savino (2014), “The information value of the stress test.” *Journal of Money, Credit and Banking*, 46. [0768]

Morris, Stephen, Daisuke Oyama, and Satoru Takahashi (2024), “Implementation via information design in binary-action supermodular games.” *Econometrica*, 92, 775–813. [0766]

Morris, Stephen and Hyun Song Shin (2006), “Global games: Theory and applications.” *Advances in Economics and Econometrics*, 56. [0780]

Morris, Stephen and Ming Yang (2022), “Coordination and continuous stochastic choice.” *The Review of Economic Studies*, 89, 2687–2722. [0767]

Myerson, Roger B. (1986), “Multistage games with communication.” *Econometrica*, 323–358. [0765]

Orlov, Dmitry, Pavel Zryumov, and Andrzej Skrzypacz (2023), “The design of macroprudential stress tests.” *The Review of Financial Studies*, 36, 4460–4501. [0767]

Petrella, Giovanni and Andrea Resti (2013), “Supervisors as information producers: Do stress tests reduce bank opaqueness?” *Journal of Banking and Finance*, 37, 5406–5420. [0768]

Quigley, Daniel and Ansgar Walther (2024), “Inside and outside information.” *The Journal of Finance*. [0768]

Rochet, Jean-Charles and X. Vives (2004), “Coordination failures and the lender of last resort: Was bagethor right after all?” *Journal of the European Economic Association*, 2, 1116–1147. [0768]

Sakovics, Jozsef and Jakub Steiner (2012), “Who matters in coordination problems?” *American Economic Review*, 102, 3439–3461. [0766]

Segal, Ilya (2003), “Coordination and discrimination in contracting with externalities: Divide and conquer?” *Journal of Economic Theory*, 113, 147–181. [0766]

Shimoji, Makoto (2021), “Bayesian persuasion in unlinked games.” *International Journal of Game Theory*, 1–31. [0767]

Szkup, Michal and Isabel Trevino (2015), “Information acquisition in global games of regime change.” *Journal of Economic Theory*, 160, 387–428. [0767]

Taneva, Ina (2019), “Information design.” *American Economic Journal: Microeconomics*, 11, 151–185. [0767]

Williams, Basil (2017), “Stress tests and bank portfolio choice.” Manuscript, Imperial College. [0767]

Winter, Eyal (2004), “Incentives and discrimination.” *American Economic Review*, 94, 764–773. [0766]

Yang, Ming (2015), “Coordination with flexible information acquisition.” *Journal of Economic Theory*, 158, 721–738. [0767]

Co-editor Simon Board handled this manuscript.

Manuscript received 7 July, 2023; final version accepted 18 August, 2024; available online 9 September, 2024.