

Keeping in the Dark with Hard Evidence
Daniel Bird, Alexander Frug

Working paper 8-2026

The Foerder Institute for Economic Research
and
The Sackler Institute of Economic Studies

Keeping in the Dark with Hard Evidence

Daniel Bird* Alexander Frug[†]

August 4, 2026

Abstract

We present a dynamic learning setting in which the periodic data observed by a decision-maker are mediated by an agent. We study when, and to what extent, this mediation can distort the decision-maker's long-run learning, even though the agent's reports are restricted to consist of verifiable hard evidence and must adhere to certain standards. We first illustrate the extent and mechanisms of manipulation in specific economic settings. We then derive a general manipulation-proof law of large numbers: when it holds, the decision-maker's learning is guaranteed in the long run; when it fails, the scope for manipulation is essentially unrestricted.

*Eitan Berglas School of Economics, Tel Aviv University (e-mail: dbird@tauex.tau.ac.il).

[†]Department of Economics and Business, Universitat Pompeu Fabra and Barcelona School of Economics (e-mail: alexander.frug@upf.edu).
Bird acknowledges financial support from the Foerder Institute for Economic Research.

1 Introduction

Classical statistics assures us that increasingly large samples of independent observations lead to arbitrarily accurate inference and correct decisions. But what if the data are filtered by someone with an agenda? In many practical environments – from employee evaluation to policy choices by government agencies – the decision-maker never sees the raw data. Instead, data are processed, summarized, and presented by self-interested agents. When agents have even limited discretion over how to present the data – e.g., what pieces of evidence to disclose, how to summarize large data into concise reports, which statistical method to apply – the statistical learning problem becomes entangled with strategic behavior.

This paper studies, in a dynamic setting, when such strategic reporting undermines the Law of Large Numbers (LLN), and just how far agency frictions can distort long-run inference. We first illustrate the strength and discontinuous nature of (limited) flexibility in periodic reporting on long-run inferences in two stylized settings. We then derive a general *Manipulation-Proof Law of Large Numbers* (MP-LLN), and analyze the feasible scope for manipulation when it fails.

Consider a “reputation management” setting in which, in each period, a decision-maker must decide whether to (temporarily) “accept” or “reject” an agent. The agent knows his type – Good or Bad – and seeks to be accepted as often as possible. Periodic information and reports work as follows: in each period, the agent observes two conditionally independent realizations of an informative binary signal and discloses exactly one of them (a structure we term *selective forced disclosure*). The decision-maker then updates her belief and accepts the agent for that period if and only if her updated belief that the agent is Good exceeds a critical threshold. This setting captures, for example, a corporate entity that provides periodic financial reports to obtain or retain an investment-grade rating from a credit rating agency.

We characterize when the decision-maker’s long-run beliefs – and, hence, the long-term acceptance rate – are susceptible to manipulation. It turns out that this occurs if and only if a simple condition holds: there exists a pair of state-dependent reporting strategies that generate the same periodic distribution of reports in both states. We term this condition *Keeping in the Dark* (KID).

When KID fails, the logic of LLN extends to the mediated learning setting and the decision-maker’s learning is (probabilistically) guaranteed. Interestingly, if KID holds, then the long-run manipulation of the decision-maker’s beliefs can be re-

markably effective. Not only can the long-term acceptance rate be increased for any interior prior belief, but we show that, unless the decision-maker’s prior belief belongs to a specific interval, she can be induced to accept the agent, *in equilibrium*, with probability arbitrarily close to the theoretical upper bound characterized by Kamenica and Gentzkow (2011).

Recall that in Kamenica and Gentzkow, achieving maximal persuasion required the agent to have access to all possible (direct) signals and commit to an experiment before learning his type. In our setting, by contrast, the same long-term acceptance rate can be approximately achieved in equilibrium, using a simple, exogenously given, binary signal.

To attain the belief manipulation described above, the agent’s reporting strategy must rely on private information. In the reputation management setting, the agent was fully informed about his type. To study the precise role of the agent’s private information in shaping the decision-maker’s long-run beliefs, we turn to a setting in which both players are initially uninformed. Hence, the agent learns from unmediated signals provided by nature over time, while the decision-maker continues to learn through the agent’s periodic reports. Moreover, to isolate the role of private information in the cleanest possible way, the setting we consider features an agent whose goal is not to steer the decision-maker’s belief in a particular direction, but simply to prevent her from learning. Beyond clarifying the role of information asymmetry and characterizing the level of private information required to block the decision-maker’s learning, we believe this setting is both novel and economically relevant, making it a natural object of study in its own right.

Consider the following “deep state” setting. A politician (decision-maker) wants to learn whether the state of the world is *Green* or *Blue* in order to enact an appropriate irreversible reform (she may also maintain the status quo). Implementing a reform is costly for a bureaucrat (agent), so he prefers to preserve the status quo. As before, every period the bureaucrat observes two conditionally independent realizations of a binary signal from which he must disclose one. However, both the politician and the bureaucrat are initially uninformed. We ask whether, despite beginning with no private information, the bureaucrat can craft the periodic reports to curtail the politician’s learning and thereby prevent reform.

The main step in the analysis fully characterizes what private information is needed to fully block the decision-maker’s learning in a given period. It turns out

that whether this degree of private information can be attained in finite time collapses to essentially the same KID condition with a mild modification: instead of merely requiring that the state-dependent report distributions intersect, S-KID requires that their intersection has a non-empty interior.

We illustrate two qualitatively distinct but representative cases in Figure 1, employing a symmetric version of the setting where the prior is uniform and each realization of the binary signal – of which the bureaucrat observes two realizations every period – matches the state with probability $\phi > \frac{1}{2}$.

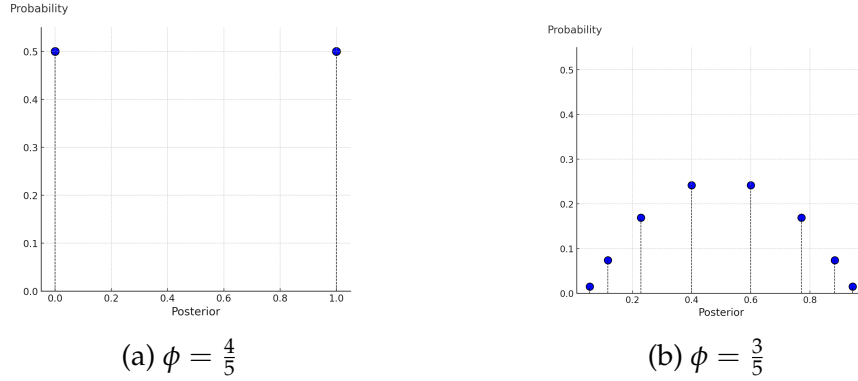


Figure 1: Distribution of the decision-maker's long-run beliefs.

Suppose first that $\phi = \frac{4}{5}$. In this case, S-KID fails and, regardless of the bureaucrat's strategy, the politician will learn the state in the long run. On the other hand, if $\phi = \frac{3}{5}$, S-KID holds and the bureaucrat can induce the long-run belief distribution depicted in Figure 1b. This can be done using a simple strategy: Initially, every period the agent reveals one of the two realizations at random. During that time, the agent's private information is insufficient to block the decision-maker's learning. It turns out that, *regardless of the disclosed signal realizations in the first 7 periods*, the required amount of private information is accumulated and from that period onward the decision-maker's learning can be blocked forever.

To interpret the belief distribution in Figure 1b within the deep-state scenario, consider the case where the politician requires a 95% confidence level to initiate a reform. In this case, under the above reporting strategy, this level of confidence is never attained so she will maintain the status quo forever. That is, the bureaucrat can prevent the reform despite the two-sided learning and the requirement to provide an infinite stream of hard-evidence about the state of nature.

While the reputation management and the deep-state scenarios differ in the play-

ers’ objectives, initial information, and even the frequency of the decision-maker’s actions, they share important common features. The simplicity of the periodic sampling and reporting, combined with concrete players’ objectives allows us to analyze equilibrium long-run beliefs and, in certain cases, to solve for the induced belief distributions explicitly.

In the final part of the paper, we turn to a broader class of mediated learning problems that abstracts away from particular players’ objectives, agent’s reporting structure, or the decision-maker’s actions. We consider a general setting in which, each period, the agent observes a random sample drawn from a state-dependent distribution and submits a report subject to sample-specific constraints. These constraints can represent a wide range of real-world partial manipulations, including cherry picking of which data to disclose, informal summaries, data cleaning, data fabrication, and more. Again, the agent has some, but limited, freedom in translating samples into reports, and the decision-maker relies on these mediated reports to learn the underlying state.

In this general mediated learning environment, we shift focus away from studying equilibrium strategies or optimal behavior under specific objectives. Instead, we ask what the feasible boundaries of long-run manipulation are. We characterize the entire set of long-run belief distributions that can be induced. To do so, we derive a structural property which we term the *Manipulation-Proof Law of Large Numbers*, which determines whether long-run learning is guaranteed regardless of the agent’s strategy. When MP-LLN holds,¹ the decision-maker’s beliefs converge to the truth over time, extending the classical logic of the LLN to strategically filtered data.

The main message of this part of the paper concerns the limits of manipulation when the MP-LLN fails. Since the decision-maker is Bayesian, her belief process must satisfy the martingale property. Hence, in particular, her long-run belief distribution must average to her prior (i.e., be Bayes-plausible). It turns out that Bayes-plausibility is essentially the only constraint. Specifically, we show that if the MP-LLN does not hold on a given subset of the state space, then, within that subset, any Bayes-plausible long-run belief distribution with finite support can be induced. In this sense, the MP-LLN serves as a sharp boundary between mediated learning environments where learning is guaranteed and those where extreme manipulation is possible.

¹This is equivalent to the failure of KID in the selective forced disclosure setting considered above.

The paper proceeds as follows. In Section 2 we formally define the selective forced disclosure model of dynamic information arrival and reports. In Sections 3 and 4 we analyze, respectively, the reputation management and reform avoidance settings. In Section 5 we study the general mediated-learning framework, and discuss the MP-LLN and its failure. In Section 6 we review the related literature. All proofs are relegated to the Appendix.

2 Selective Forced Disclosure

We now develop a simple model in which the agent has some, but limited, ability to manipulate periodic reports. We then build on this specification and consider two scenarios to study the scope of long-run manipulation in equilibrium and examine how the agent's private information shapes the decision-maker's beliefs. In Section 5, we extend the analysis to a general mediated-learning framework, in which we establish general feasibility bounds on long-run belief manipulation.

Suppose there are two states B and G and let the prior probability that the state is G be $\mu_0 \in (0, 1)$. A baseline signal

$$\langle \{b, g\}, \{\theta_\omega\}_{\omega \in \{B, G\}} \rangle$$

consists of a signal realization space $\{b, g\}$ and a conditional distribution represented by the probability $\Pr(g|\omega) = \theta_\omega$ of the signal realization being g in state $\omega \in \{B, G\}$. We assume that $\theta_B < \theta_G$, meaning that a signal realization of g is indicative of state G , while a realization of b is indicative of state B .

In each period, the agent privately observes two conditionally independent realizations of the baseline signal and must disclose exactly one of them.

This model imposes simple bounds on the distribution of the agent's periodic reports in each state. The lower bound on the probability of disclosing a signal realization of g in state ω corresponds to the agent's strategy under which he discloses a realization of b whenever possible. Under this strategy, g is disclosed only if both of the signal realizations are g , an event that occurs with probability θ_ω^2 . The upper bound on this probability is attained by the opposite strategy under which the agent discloses a realization of g whenever possible. Under this strategy, a signal realization of g is disclosed unless both the signal realizations are b , an event that

occurs with probability $1 - (1 - \theta_\omega)^2$. By randomizing between disclosing g and b when both are available, the agent can implement any disclosure probability between these two bounds.

If the minimal probability with which g must be disclosed in state G (i.e., θ_G^2) is strictly greater than the maximal probability with which it can be disclosed in state B (i.e., $1 - (1 - \theta_B)^2$), then the long-run frequencies of disclosing g in the two states are bounded away from one another. Hence, regardless of the agent's disclosure strategy, the decision-maker will eventually learn the state (we establish this formally in Proposition 4). Therefore, a necessary condition for manipulation to be feasible is:

$$\theta_G^2 \leq 2\theta_B - \theta_B^2. \quad (\text{KID})$$

When this condition is satisfied, it opens the possibility for agent types B and G to “collude” on a periodic uninformative disclosure strategy. Specifically, since the sets of distributions over disclosed evidence across the two states intersect, there exists a strategy profile under which the probability of disclosing a signal realization of g is the same in both states. Therefore, we refer to this condition as Keeping in the Dark.

3 The story of reputation management

Consider an agent who knows his type – Good (G) or Bad (B) – and seeks repeated acceptance from a decision-maker who updates her beliefs based on the hard evidence disclosed over time. The decision-maker accepts the agent in a given period if her belief that the agent is Good, in that period, meets or exceeds a given threshold $\nu^* \in (0, 1)$. Note that the decision-maker's choice to accept/reject is period-specific and can be revised every period as learning continues. The agent discounts future payoffs at rate $\beta < 1$. We conduct an equilibrium analysis in which, at every period, the agent maximizes his expected discounted continuation payoff.² Our objective is to characterize the highest long-run acceptance rate that can be sustained in equilibrium.³

²Since any finite sequence of disclosed realizations occurs with positive probability, the agent's strategy profile uniquely pins down the decision-maker's belief via Bayes' law. Hence, we do not specify the decision-maker's beliefs or consequent actions.

³The equilibrium that maximizes the long-run acceptance is not only an extreme point of the set of equilibria, but it is also an equilibrium that maximizes the agent's expected payoff under the veil of ignorance when the discount factor converges to one.

We find that, if condition KID holds, long-run persuasion can be remarkably effective. Specifically, in many cases, the agent can approximate the strongest long-run belief manipulation consistent with Bayes-plausibility, despite the exogenous constraints on the feasible reporting strategies and his inability to commit to a disclosure strategy. That is, the equilibrium long-run acceptance rate can be made arbitrarily close to the theoretical upper bound characterized by Kamenica and Gentzkow (2011). We refer to this upper bound as the Kamenica-Gentzkow Bound (KGB).

3.1 Mechanisms of reputation management: freeze, mix & wait

To describe a possible agent's strategy for manipulating the long-run beliefs of the decision-maker, we take as a starting point the following naive strategy: every period, if at least one of the signal realizations is g , the agent discloses g . We term this strategy *within-period sanitization*.⁴ If condition KID does not hold, the disclosure of a signal realization of g increases the decision-maker's belief for any disclosure strategy. Therefore, in such cases, within-period sanitization is the unique equilibrium and the decision-maker will eventually learn the agent's type. Hence, in the subsequent analysis we assume that condition KID holds.

A natural idea for boosting the long-run acceptance rate above what can be attained via within-period sanitization is to use the *early sanitization* strategy: so long as the decision-maker's period t belief is outside the acceptance region (i.e., is less than ν^*), the agent plays within-period sanitization, and once the belief enters the acceptance region the agent switches to playing an uninformative disclosure strategy profile (the existence of which follows from condition KID) thereby freezing the decision-maker's beliefs.

This strategy is simple and effective: it replaces full learning in the limit with a distribution of long-run beliefs that includes 0 (reflecting the decision-maker's certainty that the agent is of type B) and beliefs within the interval $[\nu^*, 1)$. Moreover, this strategy and the corresponding decision-maker's beliefs constitute an equilibrium.⁵

⁴This terminology follows the one introduced by Song Shin (2003), who analyses voluntary disclosure of good/bad evidence by firms.

⁵To see that early sanitization is an equilibrium, note that the decision-maker's beliefs are not altered by periodic disclosures once the agent enters the acceptance region, and disclosing g before this occurs increases the decision-maker's belief. Therefore, disclosing g leads to a quicker entrance into the acceptance region.

Clearly, if the decision-maker's prior belief lies within the acceptance region, $\mu_0 \geq v^*$, the early sanitization strategy trivially achieves maximal acceptance, as the agent is accepted in every period. If $\mu_0 < v^*$, early sanitization still increases the long-run acceptance rate relative to full learning by the decision-maker, but it typically leads to overshooting: beliefs cross v^* at points strictly above v^* , and so the resulting distribution of long-term beliefs is bounded away from the KGB.⁶ To further increase the long-term acceptance rate, the agent must exert finer control over the belief path.

To define more refined manipulation strategies, let $v_{-1}^* \in (0, v^*)$ denote the belief at which, if the agent plays within-period sanitization and discloses g , the decision-maker will update her belief to exactly v^* (the acceptance threshold). Similarly, let $v_{-2}^* \in (0, v_{-1}^*)$ denote the belief at which, if the agent plays within-period sanitization and discloses g , the decision-maker will update her belief to exactly v_{-1}^* .

Consider first the case where $\mu_0 \in (0, v_{-1}^*)$. We now illustrate a modification of the early sanitization strategy that will serve as a key building block in constructing an equilibrium strategy profile that approximates the KGB.

1-mixing early-sanitization strategy: The agent plays early sanitization until the decision-maker's belief enters (v_{-2}^*, v_{-1}^*) for the first time. When this occurs, the agent's strategy depends on his type. Type B continues to play within-period sanitization (for one period). In contrast, type G – if he samples evidence $\{bg\}$ – mixes over which signal realization to disclose. Specifically, he discloses g with probability α , chosen so that upon observing g , the decision-maker's updated belief is exactly v_{-1}^* . The existence of such α follows from condition KID.

For this mixing to be part of an equilibrium, type G must be indifferent between disclosing g and b . Since the choice of α ensures that the disclosure of g raises the decision-maker's belief, it follows that a disclosure of b must lower it. Hence, to obtain such indifference, if g is disclosed, the interaction enters a waiting phase of τ (possibly stochastic) periods, during which the agent types play an uninformative periodic disclosure strategy profile.⁷ After the waiting phase ends, the agent resumes early

⁶In the KGB, the prior belief $\mu_0 \in (0, v^*)$ is replaced with a Bayes-plausible distribution over two posterior beliefs: 0 and v^* .

⁷Implementing a stochastic waiting time requires a public randomization device or the ability to

sanitization (if b is disclosed in the mixing period, early sanitization resumes immediately).

The 1-mixing early sanitization strategy uses a periodic uninformative disclosure strategy as an incentivization device. Specifically, the temporary waiting phase introduces a cost of delay: during this phase, the decision-maker's belief is outside the acceptance region and the agent receives his lowest periodic payoff. Crucially, as the decision-maker perceives any disclosure as uninformative during this phase, the agent cannot bypass the waiting phase to accelerate acceptance.⁸ The duration of this phase is calibrated so that the agent's expected discounting until entry to the acceptance region does not depend on which signal is disclosed in the mixing period.

To approximate the maximal Bayes-plausible acceptance rate in the long run, we extend the idea described above to allow for more than one mixing period. Intuitively, suppose that after the first mixing period the belief falls below ν_{-1}^* . Instead of continuing with the early sanitization strategy (like in the 1-mixing early-sanitization strategy), we can apply the logic of the 1-mixing strategy again. From the ex-ante perspective, this would then allow for two mixing periods, and the number of mixing periods can be extended further. By allowing for multiple potential mixing periods, we show that the probability of entering the acceptance region at ν^* can be made arbitrarily close to the overall probability of eventual acceptance.

Of course, the proof of Proposition 1 below requires additional work. For example, anticipating the possibility of future mixing periods affects the distribution of the time to acceptance following earlier mixing periods, which in turn alters the duration of the waiting phases throughout. Nonetheless, the high-level intuition for this result is captured by the structure described above.

Proposition 1. *Assume that Condition KID holds. The maximal Bayes-plausible acceptance rate can be approximated, whenever $\mu_0 \leq \nu_{-1}^*$.*

An interesting insight from Proposition 1 is that to attain maximal persuasion

engage in a two-sided cheap-talk communication to create an appropriate jointly controlled lottery á la Aumann and Hart (2003) and Krishna and Morgan (2004). As we show below, we only require simple lotteries with two potential outcomes but, nevertheless, to facilitate exposition we avoid the formal construction of such lotteries.

⁸The idea of using costly waiting times to influence behavior has already appeared in a number of other settings (see, e.g., Escobar and Zhang, 2021, Antler, Bird and Oliveros, 2023, and Eliaz, Fershtman and Frug, 2024).

in equilibrium, it may be necessary to employ a periodic uninformative disclosure strategy profile not only after, but also before, the decision-maker’s belief enters the acceptance region. Once the belief is already in that region, making the periodic disclosure uninformative allows the agent to protect himself from ever being rejected. In contrast, applying this strategy profile prior to entering the acceptance region helps shape equilibrium incentives and increase the overall acceptance rate in the long run.

The proof of Proposition 1 is constructive: we derive a specific strategy and analyze its properties. However, this construction does not work for the case where $\mu_0 \in (\nu_{-1}^*, \nu^*)$, as then we cannot impose a waiting cost on the agent to incentivize mixing in the first period of the interaction. It turns out that, in this case, the maximal Bayes-plausible acceptance rate cannot be approximated in equilibrium. We show this result formally in Appendix B.

4 The story of reform avoidance: a model of deep state

To achieve maximal long-run persuasion in the scenario analyzed in Section 3, the different agent-types engaged in a delicate and coordinated manipulation of the decision-maker’s beliefs. Importantly, the prescribed agent’s strategy relied on him *knowing* the state of the world. In many situations, however, the agent may not have such information. He might begin the interaction with only partial knowledge – perhaps as uninformed as the decision-maker, slightly more informed, or even less. Is the agent’s private information necessary for manipulation? And if so, “how much” private information is needed?

To address these questions, we now analyze a setting where the initially uninformed agent gradually learns from periodic samples provided by nature, while the decision-maker learns from the agent’s periodic reports. For tractability and clean parametric analysis, we employ the same selective forced disclosure evidence structure that we defined in Section 2 and used in Section 3. Moreover, building on the observation from the previous section – that a necessary condition for long-run belief manipulation is the ability to generate indistinguishable report distributions across states – we now consider a setting in which preventing the decision-maker’s learning is the agent’s *objective*, rather than a means to some other end. In particular, we develop and analyze a simple model of a “deep state.”

Consider an interaction between a newly elected politician (decision-maker), who seeks to enact a reform – provided she becomes sufficiently convinced about which reform is appropriate – and a bureaucrat (agent), who must provide periodic reports and implement whatever reform the decision-maker decides to pursue. Implementing a reform is costly for the bureaucrat, and so he prefers to preserve the status quo. Thus, he has an incentive to prevent the decision-maker’s learning.⁹

Formally, the state is either *Blue* or *Green*, and at the start of the interaction, the players share a common prior belief μ_0 that the state is *Green*. The politician can make an irreversible decision – ending the interaction – to enact one of two potential reforms, $\mathcal{B}$ or $\mathcal{G}$. She enacts reform $\mathcal{B}$ (respectively, $\mathcal{G}$) the moment her belief that the state is *Blue* (*Green*) exceeds the threshold $1 - \pi$, for some $\pi < 1/2$. As long as her confidence in the realized state remains below $1 - \pi$, she maintains the status quo and the interaction continues.

The impact of the agent’s private information – or lack thereof – is most noticeable in the first period of the interaction. If the agent’s disclosable evidence is $\{bb\}$ or $\{gg\}$, he has no choice over what to disclose. To make the overall disclosure uninformative, the agent would need to offset the state-dependence induced by these forced cases by strategically choosing what to disclose when his available evidence is $\{bg\}$. However, upon sampling $\{bg\}$, the agent’s belief is the same in both states. Hence, the choice of what to disclose cannot correlate with the state. To see that the overall disclosure in period 1 must be informative, note that the agent’s strategy in that period can be fully summarized by a single parameter $\alpha \in [0, 1]$, the probability of disclosing g when the realized evidence is $\{bg\}$. Letting $d \in \{b, g\}$ denote the disclosed signal, in state ω we have

$$\begin{aligned} \Pr(d = g \mid \omega) &= \Pr(gg \mid \omega) + \alpha \Pr(bg \mid \omega) \\ &= \theta_\omega^2 + 2\alpha \theta_\omega(1 - \theta_\omega) \\ &= (1 - \alpha)\theta_\omega^2 + \alpha(2\theta_\omega - \theta_\omega^2). \end{aligned}$$

Since $\theta_G > \theta_B$, this expression is strictly larger in state G than in state B for every $\alpha \in [0, 1]$ (as a convex combination of θ_ω^2 and $2\theta_\omega - \theta_\omega^2$, both of which are strictly

⁹Arieli et al. (2024) also study a model where the sender tries to persuade a receiver while (gradually) learning about the state of nature. In contrast to us, they consider a setting where in each period the sender can generate *any* signal that is measurable with respect to her periodic information, and assume that her goal is to persuade the receiver that the state is high.

increasing in θ_ω on $[0, 1]$). Hence, $\Pr(d = g \mid G) > \Pr(d = g \mid B)$ under any feasible strategy, and so the disclosed signal is necessarily informative in period 1.

Put differently, an uninformed agent has only one “degree of freedom” (how to behave on $\{bg\}$), and because that behavior cannot depend on the state, it cannot offset the informational content coming from the forced cases $\{gg\}$ and $\{bb\}$. To render disclosure uninformative, the agent needs additional private information that will generate heterogeneity among $\{bg\}$ -types and allow for a state-dependent behavior that the decision-maker cannot invert.

Since every period the agent observes two conditionally independent signal realizations but discloses only one (according to whatever rule), the agent accumulates private information over time. Intuitively, as time goes by, the amount of private information grows and may eventually allow for a within-period uninformative disclosure strategy. The main focus of this section is to determine when it is possible for the agent to stop the decision-maker’s learning, at least from some period onwards.

We begin our analysis by characterizing the agent’s (private) information structures that enable him to make the periodic disclosure uninformative. We will later apply this characterization in analyzing whether the bureaucrat can prevent the reform.

4.1 One-shot forced disclosure with hard and soft information

Consider a single-period forced disclosure game in which prior to collecting and disclosing evidence, the agent privately observes an undisclosable signal (soft information),

$$\langle \mathcal{S}, \{\Pr(\cdot|\omega)\}_{\omega \in \{B, G\}} \rangle,$$

where $\mathcal{S}$ is a finite realization space and $\Pr(\cdot|\omega)$ is a probability distribution over $\mathcal{S}$ given the state ω . Next, as before, the agent samples two realizations from the baseline signal, of which he must disclose exactly one. We assume that the soft information and baseline signal realizations are conditionally independent.

The agent’s type-space at the time of selecting what evidence to disclose is

$$Y = \{\{bb\}, \{bg\}, \{gg\}\} \times \mathcal{S},$$

where the first component is the sampled evidence from which he can disclose exactly one realization and the second is soft information. For each $y \in Y$, we denote

by $\mu(y)$ the agent's updated belief based on both components of his information. To state our characterization we first define a specific disclosure strategy and an intuitive notion of the inferences made from disclosure.

Definition (Belief-contrarian strategy). *We say that the agent uses a belief-contrarian strategy if an agent-type with evidence $\{bg\}$ discloses g if $\mu(y) < \mu_0$ and discloses b if $\mu(y) > \mu_0$.*

Roughly speaking, an agent that uses such a strategy attempts to mislead the decision-maker by disclosing the evidence that goes against the overall information that he has received.¹⁰

Definition (Meaning reversal strategy). *We say that the agent's strategy leads to weak meaning reversal if disclosing g weakly decreases the decision-maker's belief that the state is G .*

The next result characterizes the agent's private information structures under which an uninformative disclosure strategy profile exists.

Proposition 2. *The belief-contrarian strategy leads to weak meaning reversal if and only if there exists an uninformative disclosure strategy profile.*

Proposition 2 gives rise to a simple condition that determines whether the agent has sufficient private information to render the periodic disclosure uninformative. To present this condition formally, let

$$D \equiv \{s \in \mathcal{S} : \mu(\langle bg, s \rangle) \leq \mu_0\}.$$

That is, D is the set of agent-types who can disclose either b or g and have updated the belief that the state is G weakly downward after they receive their information.

Claim 1. *The belief-contrarian strategy results in weak meaning reversal if and only if:*

$$\Pr(gg|G) - \Pr(gg|B) \leq \Pr(bg, D|B) - \Pr(bg, D|G). \quad (1)$$

Moreover, if (1) holds for some $\mu_0 \in (0, 1)$, then it holds for all such μ_0 .

¹⁰the disclosure conditional on $\mu(y) = \mu_0$ is immaterial.

Claim 1 characterizes the existence of an uninformative disclosure profile by inequality (1). The left-hand side of (1) is prior-free. On the right-hand side, the only apparent dependence on the prior μ_0 is through D . However, for any given realization s , the inequality $\mu(\langle bg, s \rangle) \leq \mu_0$ is equivalent to a likelihood-ratio test that does not involve μ_0 :

$$\mu(\langle bg, s \rangle) \leq \mu_0 \iff \frac{\Pr(bg, s \mid G)}{\Pr(bg, s \mid B)} \leq 1.$$

Hence, membership in D is independent of μ_0 , so if an uninformative disclosure profile exists for some $\mu_0 \in (0, 1)$, then it exists for all $\mu_0 \in (0, 1)$. Another implication of this prior-independence is that allowing the decision-maker to have private information beyond the prior would not affect the existence of an uninformative disclosure strategy profile. In particular, the agent and the decision-maker could privately observe different signals, yet the feasibility of uninformative disclosure is driven entirely by how much private information the agent has.

4.2 From one-shot to dynamics: reform avoidance revisited

Building on the condition given in (1), we now return to the dynamic setting. Suppose that the agent's disclosure strategy in a given period prescribes revealing each of the realizations with probability $\frac{1}{2}$. This has the following implications.

First, the disclosed realization updates the decision-maker's belief as if she had observed a single draw of the baseline signal. Second, the agent's informational advantage can be represented by the number of privately observed draws from the baseline signal. This is the case because the disclosed and concealed signal realizations are conditionally independent – a property that is preserved not only since the draws are conditionally independent but also through the agent's disclosure strategy. Third, if there is a disclosure history of length t following which the agent can block the decision-maker's future learning, then due to the prior independence property of the condition in (1) he can do so for any history of disclosed realizations of length t .

It follows that for this disclosure strategy, the task of checking whether the required informational advantage can eventually be attained is reduced to checking whether privately observing sufficiently many draws of the baseline signal – or equivalently, whether being arbitrarily close to being fully informed – provides the required informational advantage. The next result shows that this requirement is

equivalent to

$$\theta_G^2 < 2\theta_B - \theta_B^2. \quad (\text{S-KID})$$

That is, condition S-KID is equivalent to condition KID holding with strict inequality, so that the intersection of the feasible state-dependent report distributions has a non-empty interior.

Proposition 3. *There exists $T^* \in \mathbb{N}$ such that the agent can prevent the decision-maker from learning any information from time T^* onward if and only if condition S-KID holds. Moreover, when such T^* exists it is independent of the prior.*

Although Proposition 3 may appear asymptotic, this is not the case. To illustrate this, consider a symmetric version of the selective forced disclosure model in which both signals have the same precision. That is, assume that $\phi \equiv \theta_G = 1 - \theta_B$. Note that in this symmetric model we require that $\phi > 1/2$ to make signal realization g indicative of state G , and that condition S-KID becomes $\phi < \frac{1}{\sqrt{2}}$. Figure 2 depicts the T^* that corresponds to the initial random disclosure strategy as a function of ϕ and shows that, for most of the relevant range of ϕ , the decision-maker's learning can be stopped after a relatively short amount of time.

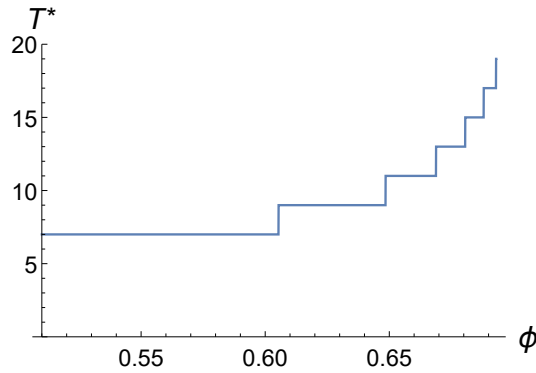


Figure 2: T^* as a function of ϕ .

Recall the example from the Introduction in which $\phi = \frac{3}{5}$. As can be seen in Figure 2 the informational advantage required to block the decision-maker's learning for $\phi = \frac{3}{5}$ can be attained after 7 periods in which the agent discloses at random. During this phase, the decision-maker observes 7 informative signals (under the random disclosure strategy the order in which these signals are disclosed is immaterial), which leads to the 8 possible beliefs presented in Figure 1b.

Whether this simple strategy is optimal depends on the decision-maker’s policies. In that example, the induced decision-maker’s beliefs after 7 periods belong to the interval $[0.055, 0.945]$. Thus, if the decision-maker enacts the reform only when she is certain about the state with at least 95%, this simple strategy prevents the reform regardless of which signal realizations are disclosed in these 7 periods.

If the decision-maker is more lenient, the agent may still be able to prevent a reform by using a more complicated strategy. Under random disclosure the disclosed signal realization is uninformative about the undisclosed one. However, this strategy does not attempt to prevent the decision-maker from developing extreme beliefs while the required informational advantage is obtained. More complex strategies that consider the later objective can prevent the reform even with a more lenient decision-maker.¹¹ Characterizing the optimal reform avoidance strategy is beyond the scope of this paper.

5 Manipulation in general mediated learning

Up to this point we have focused on the selective forced disclosure framework, applied to concrete environments, to identify forms of long-term manipulation that may arise in equilibrium. The potential for long-run manipulation in mediated learning is clearly not confined to these cases; it exists in any setting in which an agent has limited, but nontrivial, discretion over how to present data to a decision-maker.

In this section, we consider a general mediated learning setting. We introduce the Manipulation-Proof Law of Large Numbers (MP-LLN), which specifies when the decision-maker’s long-run learning is unaffected by strategic reporting.¹² We then study the limits of manipulation when the MP-LLN fails. Rather than analyzing equilibrium outcomes in a particular setting, we investigate the extent to which the decision-maker’s learning can be manipulated, as determined directly by the structure of the periodic data available to the agent and the reporting protocol that ties his reports to the privately observed samples. To characterize the boundaries

¹¹For example, it can be shown that, for $\phi = 3/5$, the agent can prevent the reform even if the decision-maker enacts a reform when her certainty about the state exceeds 92.5%.

¹²When the MP-LLN holds, even though the decision-maker will learn the true state, the agent may be able to impact the rate at which the decision-maker learns. In this paper we set aside this aspect of dynamic mediated learning.

of manipulation, we consider the case in which the agent is fully informed about the state of the world. We show that the scope of long-run manipulation is essentially bounded only by Bayes-plausibility, restricted to subsets of states that satisfy a simple joint manipulability condition introduced below.

Framework

Consider a dynamic mediated learning problem in which the state of nature $\omega \in \Omega$ is distributed according to a prior distribution μ_0 with full support. Information about the state is generated by a sequence of periodic samples $\{x_t\}_{t=1}^T$ that are drawn i.i.d. according to a family of distinct state-dependent distributions $\{\chi_\omega\}_{\omega \in \Omega}$. Let X be the set of possible periodic sample realizations. We assume that both Ω and X are finite.

Every period the agent privately observes the realized sample $x \in X$, and submits a report to the decision-maker. The key notion that captures the agent's flexibility in manipulating the raw data is the feasible report correspondence $R(x)$ that specifies the set of feasible reports for each periodic sample $x \in X$.¹³ We assume that the set of all feasible reports, $R = \bigcup_{x \in X} R(x)$, is finite. The agent's periodic strategy is a mapping $\sigma : X \rightarrow \Delta(R)$ such that $\sigma(x) \in \Delta(R(x))$, which, for every realization of the periodic sample, specifies a probability distribution over feasible reports. The agent's choice of periodic strategy in a given period can depend on his information and, in particular, on the state of nature.

This framework captures a wide array of phenomena (e.g., cherry picking, partial data fabrication, and strategic cleaning of the data) in which the agent's reports are grounded in real data, but typically do not reveal the full truth. For example, the selective forced disclosure framework described in Section 2 can be embedded in our general framework by setting $X = \{bb, bg, gg\}$, $R = \{b, g\}$ and

$$R(bb) = \{b\}, R(gg) = \{g\}, R(bg) = \{b, g\}.$$

A simple model of occasional data-fabrication where the agent observes one signal, but can sometimes fabricate its realization before disclosing it can be embedded as follows. The periodic sample is $\langle s, f \rangle$, where $s \in S$ is the realization of the signal,

¹³This object is akin to the feasible reporting set in Dye (1988) or the set of statements in Glazer and Rubinstein (2006).

and $f \in \{0, 1\}$ is an indicator for whether or not the agent can fabricate the signal realization. The set of reports is $R = S$, and

$$R(s, 0) = \{s\}, R(s, 1) = R.$$

We illustrate an additional example of strategic cleaning of the data at the end of this section.

5.1 The Manipulation-Proof Law of Large Numbers

In this class of mediated learning problems the decision-maker attempts to infer the state of nature from the sequence of the agent's reports $\{\hat{\sigma}_t(x_t)\}_{t=1}^T$, where $\hat{\sigma}_t(\cdot)$ is the report submitted in period t under reporting strategy $\sigma_t(\cdot)$, rather than from the sequence of periodic samples $\{x_t\}_{t=1}^T$. We now establish a condition under which the decision-maker's long-term learning is asymptotically equivalent in both cases. It turns out that even though the agent can use complex dynamic strategies, to check whether or not long-term manipulation of the decision-maker's beliefs is feasible it suffices to look at periodic strategies that do not change over time.

Given a (within-period) reporting strategy σ and a state ω , let

$$f(\cdot \mid \omega, \sigma) = \sum_{x \in X} \sigma(\cdot \mid x) \chi_\omega(x),$$

denote the induced distribution of reports in state ω under σ . Note that if the agent uses the strategy σ in every period, then, by the law of large numbers, the long-run frequency of report r in state ω is $f(r \mid \omega, \sigma)$. Denote the set of feasible report distributions in state $\omega \in \Omega$ by

$$\mathcal{F}(\omega) = \left\{ f(\cdot \mid \omega, \sigma) : \sigma \text{ is a within-period reporting strategy} \right\}.$$

Note that these sets are closed and convex.

The key question in determining if the decision-maker will be able to learn the state regardless of the agent's strategy is whether the sets $\{\mathcal{F}(\omega)\}_{\omega \in \Omega}$ overlap. Let $\mu(\cdot)$ denote the decision-maker's updated belief, and δ_ω the degenerate belief assigning probability one to state ω .

Proposition 4 (Manipulation-Proof Law of Large Numbers). *If the sets $\{\mathcal{F}(\omega)\}_{\omega \in \Omega}$*

are pairwise disjoint, then for any agent's periodic strategy $\{\sigma_t\}_{t=1}^\infty$ that may depend arbitrarily on the past and the state of nature,

$$\mu(\{\hat{\sigma}_t(x_t)\}_{t=1}^T) \xrightarrow[T \rightarrow \infty]{D} \delta_{\omega'},$$

where ω' is the realized state of the world.

Intuitively, even if the agent distorts which reports are sent in each period arbitrarily, the logic of the law of large numbers (or more precisely the strong law of martingale differences) ensures that every converging sub-sequence of long-run reporting frequencies in state ω converges to an element of the set $\mathcal{F}(\omega)$. Since these sets are also closed and disjoint, the long-run reporting frequencies in different states can be separated, so the decision-maker can eventually identify the true state.

5.2 The scope of long-run manipulation

To study the scope of manipulation when the sets $\{\mathcal{F}(\omega)\}_{\omega \in \Omega}$ are not pairwise disjoint, we begin with the following definition.

We say that a subset $S \subseteq \Omega$ with $|S| > 1$ is *strictly manipulable* if

$$|\bigcap_{\omega \in S} \mathcal{F}(\omega)| > 1.$$

Since the sets $\mathcal{F}(\omega)$ are convex, if S is strictly manipulable, there is a continuum of reporting distributions that can be induced in all the states in S .¹⁴ Note that in the selective forced disclosure model introduced in Section 2, condition S-KID is equivalent to the set Ω being strictly manipulable.

To describe the scope of manipulation on strictly manipulable sets, we need to formalize the notion of conditional Bayes-plausibility. Given a (full-support) prior μ_0 on Ω and a subset $S \subseteq \Omega$, let

$$\mu_0^S(\omega) = \frac{\mu_0(\omega)}{\mu_0(S)} \quad \text{for } \omega \in S$$

be the prior distribution restricted to S . A finite-support distribution of beliefs

¹⁴We discuss the nongeneric case where $|\bigcap_{\omega \in S} \mathcal{F}(\omega)| = 1$ in Appendix C.

$\{(v_k, p_k)\}_{k=1}^K$ with $v_k \in \Delta(S)$ is said to be *Bayes-plausible conditional on S* if

$$\sum_{k=1}^K p_k v_k(\omega) = \mu_0^S(\omega) \quad \text{for all } \omega \in S.$$

In words, Bayes-plausibility conditional on S requires that the proposed distribution of posteriors averages back to the prior beliefs conditional on S .

Proposition 5. *Let $S \subseteq \Omega$ be a strictly manipulable set and let $\{(v_k, p_k)\}_{k=1}^K$ be a finite-support Bayes-plausible distribution of posteriors conditional on S . There exists a reporting strategy such that:*

1. *The decision-maker will identify almost surely whether or not $\omega \in S$.*
2. *If the true state lies in S , then the decision-maker's long-run belief converges to v_k with probability p_k , for each $k = 1, \dots, K$.*

Proposition 5 shows that within a strictly manipulable set the only restriction on the agent's ability to manipulate the decision-maker is Bayes-plausibility. That is, the fact that the agent crafts reports (based on periodic samples generated by a known exogenous process) in accordance with a predetermined reporting rule, does not inherently limit his ability to manipulate the decision-maker's long-term beliefs.

The intuition behind this result is straightforward. For states in S , the fact that $\bigcap_{\omega \in S} \mathcal{F}(\omega)$ contains a continuum of report distributions provides enough flexibility to vary the induced long-run frequencies without interfering with those induced in states outside of S . By randomizing across different points in this set in a specific state-dependent manner, the agent can generate in the limit any finite-support distribution over posteriors that is Bayes-plausible conditional on S .

A direct implication of Proposition 5 is that if Ω itself is strictly manipulable, then the scope of manipulation is essentially unrestricted.

Corollary 1 (Global manipulability). *Suppose that Ω is a strictly manipulable set. Then in the long run the agent can induce any finite-support Bayes-plausible distribution of the decision-maker's beliefs.*

We conclude this section by analyzing two examples that illustrate more subtle aspects of the notion of strict manipulability.

Example 1: Selective forced disclosure with three states

Consider a selective forced disclosure setting in which the agent observes two i.i.d. binary signal realizations each period (interpreted as good or bad evidence) and is required to disclose exactly one of them. For a state with probability θ of drawing good evidence, the set of inducible long-run frequencies of reporting good evidence is

$$\mathcal{F}(\theta) = [\theta^2, 1 - (1 - \theta)^2].$$

Suppose that there are three states, $\Omega = \{\text{Low}, \text{Medium}, \text{High}\}$, with probabilities of good evidence $\theta_L = 0.2$, $\theta_M \in (0.2, 0.8)$, and $\theta_H = 0.8$, respectively. In this case,

$$\mathcal{F}(L) = [0.04, 0.36], \quad \mathcal{F}(M) = [\theta_M^2, 1 - (1 - \theta_M)^2], \quad \mathcal{F}(H) = [0.64, 0.96].$$

Since $\mathcal{F}(L) \cap \mathcal{F}(H) = \emptyset$, neither the set $\{\text{Low}, \text{High}\}$ nor the set Ω is strictly manipulable. By contrast, $\{\text{Low}, \text{Medium}\}$ is a strictly manipulable set whenever $0.2 < \theta_M < 0.6$, as in this range the intervals $\mathcal{F}(L)$ and $\mathcal{F}(M)$ overlap with a non-empty interior. Similarly, $\{\text{Medium}, \text{High}\}$ is a strictly manipulable set whenever $0.4 < \theta_M < 0.8$, as, for such values, $\mathcal{F}(M)$ and $\mathcal{F}(H)$ have overlapping interiors.

Thus, for intermediate values of θ_M , both $\{\text{Low}, \text{Medium}\}$ and $\{\text{Medium}, \text{High}\}$ are strictly manipulable sets, while for extreme values of θ_M only one of these sets is. In other words, depending on the position of the Medium state, the agent can sometimes pool it with Low, sometimes with High, and sometimes with either.

Another feature that can be illustrated in this example is that, if state M is sufficiently likely, the agent can induce a distribution of the decision-maker's beliefs that will always assign sufficiently high probability to state M , even though Ω is not strictly manipulable. The significance of this observation is most apparent in the context of the reform-avoidance example where the agent's objective is to render the decision-maker's belief in a region that avoids active actions at extreme states.¹⁵

Consider $\theta_M \in [0.4, 0.6]$ (i.e., values of θ_M for which both $\{\text{Low}, \text{Medium}\}$ and $\{\text{Medium}, \text{High}\}$ are strictly manipulable). Thus, in state Low (High) the agent can induce a long-term reporting frequency of f_L (f_H) that can also be induced in state Medium by strategies σ_L (σ_H). In state Medium, at period 0, the agent can randomly choose between σ_L and σ_H with equal probability, and stick to the chosen strategy

¹⁵In brief, this can be achieved by a randomization akin to that used in the Proposition 5.

in all periods. Doing so will lead to a long-term belief of

$$\Pr(\omega = M|f_i) = \frac{1/2 \Pr(\omega = M)}{\Pr(\omega = i) + 1/2 \Pr(\omega = M)}.$$

Thus, if state Medium is sufficiently likely under the prior and the politician requires sufficient certainty about the state to enact a reform, the agent can prevent the reform even though Ω is not strictly manipulable.

Example 2: Cleaning the data

The role of an agent that collects data is often to summarize the data into a single summary statistic. However, such agents often have the flexibility to “clean the data,” i.e., remove a small number of “non-representative” observations from the sample, before calculating the required statistic. To illustrate how such a setting can be embedded in our framework, assume that there are two states Negative and Positive, and that the agent draws a sample of three conditionally independent draws of a random variable with support $\{-2, 0, 2\}$ according to the following distribution function,

States \ Outcomes	-2	0	2
Negative	$\frac{1}{3} + \zeta$	$\frac{1}{3}$	$\frac{1}{3} - \zeta$
Positive	$\frac{1}{3} - \zeta$	$\frac{1}{3}$	$\frac{1}{3} + \zeta$

where $\zeta \in (0, \frac{1}{3})$.

The agent must report the realized average of the periodic sample at every period; however, before calculating the average the agent excludes one of the three realizations. In this example, $X = \{-2, 0, 2\}^3$, $R = \{-2, -1, 0, 1, 2\}$, and

$$r \in R(x) \Leftrightarrow \exists i, j \in \{1, 2, 3\} \text{ s.t. } i \neq j \text{ and } x_i + x_j = 2r.$$

Regardless of the agent’s strategy, the sign of the reported average in a given period corresponds to the state with probability bounded away from zero. Now, take $\zeta \rightarrow \frac{1}{3}$ and suppose for a moment that the state is Positive. In this case, the probability of a negative report in a given period approaches zero. Therefore, for sufficiently large ζ , the sign of the average report in the long run will match the state with arbitrarily high probability, implying that the set $\{\text{Positive}, \text{Negative}\}$ is

not strictly manipulable.

By always excluding the highest of the three realizations in the Positive state, the agent reduces the periodic report as much as possible. Due to the symmetric structure of the example, excluding the lowest realization in the Negative state pushes the reports upwards at exactly the same rate. Accordingly, there exists a threshold $\zeta_0 \approx 0.141$ at which, under this strategy, the expected periodic reports in both states coincide, and are thus equal to zero.

However, if $\zeta = \zeta_0$, the set {Positive, Negative} is still not strictly manipulable. To see this, it is sufficient to note that even though, for $\zeta = \zeta_0$, the strategy described above generates the same expected report in both states (and the same variance of reports), yet the distribution of reports is different between the two states:

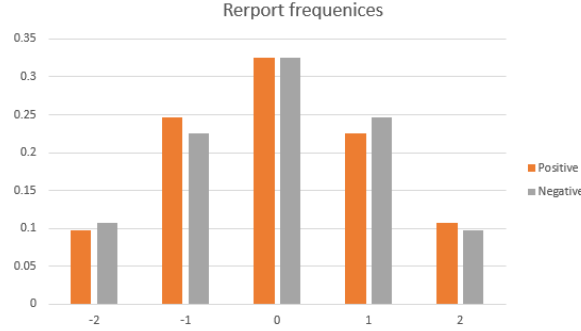


Figure 3: Reporting frequencies for extreme exclusion strategies.

The attentive decision-maker will, therefore, not be misled by the uninformative average of the periodic reports, but will infer the state based on the higher moments of the empirical frequencies of reports.¹⁶

The above discussion implies that for the set {Positive, Negative} to be strictly manipulable, it must be the case that $\zeta < \zeta_0$, which, in turn, implies that we must consider reporting strategies that are more nuanced than simply omitting extreme draws. Due to the symmetric nature of this example, generating a report distribution that is symmetric around $r = 0$ in one state is equivalent to finding a reporting distribution that can be induced in both states. For example, when $\zeta = \frac{1}{9}$, a symmetric distribution of reports is attained in the Positive state by the following reporting strategy:

¹⁶Were the agent to use any other strategy, the decision-maker would learn the state from the long-term average of the reports.

For any of the following realized samples, $\{-2, -2, 0\}$, $\{-2, 0, 0\}$, $\{-2, -2, 2\}$, remove the highest realization with probability $\frac{2}{3}$ (and the lowest realization with probability $\frac{1}{3}$); if the realized sample is $\{-2, 0, 2\}$, remove the highest realization with probability $\frac{5}{6}$ and the lowest one with probability $\frac{1}{6}$; for any other realized sample, remove the highest realization.

Employing such complex mixed strategies is necessary because achieving a perfectly symmetric distribution of reports requires simultaneous and careful fine-tuning of different regions of the report distribution.

6 Related Literature

Research on strategic disclosure of hard evidence has largely focused on static environments, where an agent seeks to prove he is of high quality. The seminal works of Grossman and Hart (1980); Grossman (1981); Milgrom (1981) establish an “unraveling” result: if disclosure is unconstrained, all private information is revealed in equilibrium, driven by the incentive of the best type to separate from lower types. Later contributions impose various constraints on disclosure. For instance, Verrecchia (1983) studies costly disclosure, while Dye (1985) and Jung and Kwon (1988) incorporate uncertainty about the sender’s information endowment.

Closer to our approach is Fishman and Hagerty (1990), who consider a static game in which an agent, observing multiple signals, must disclose exactly one from a restricted set. They examine how varying the agent’s discretion over which signals can be reported (i.e., changing the size of the set of the signals that can be disclosed) affects social efficiency. By contrast, our focus is on a dynamic environment, in which the agent can select which evidence to present over time, in order to circumvent the decision-maker’s learning via the logic of the Law of Large Numbers.¹⁷

A related literature explores constrained persuasion and selective disclosure. Glazer and Rubinstein (2004, 2006) focus on characterizing the mechanism that minimizes the probability that the listener accepts a sender she should reject in a one-shot game. More recently, Antic and Chakraborty (2024) study a static persuasion model in which a sender must choose a subset of verifiable facts to disclose and

¹⁷A related strand in accounting and finance, often referred to as Earnings Management (See Healy and Wahlen, 1999; Beyer et al., 2010 for extensive reviews of this literature), explores how managers selectively report information under various rules. While that literature shares the theme of selective disclosure, its typical assumptions differ from ours in three key respects: (i) it rarely incorporates long run learning about a fixed underlying state, (ii) managers typically aim to maximize beliefs about firm value, and (iii) report manipulation is generally costly.

characterize the optimal disclosure rule via a maximum-weight matching representation. Antic, Chakraborty and Harbaugh (2024) examine how communication between informed parties can be structured to share information while concealing it from an external monitor. Finally, Di Tillio, Ottaviani and Sørensen (2021) analyze how sub-sampling affects the informativeness of the produced evidence, contrasting random and selective sampling. These papers form the closest antecedents to our framework: they study how verifiable evidence can be strategically filtered. Our contribution differs in that we move beyond a one-shot perspective and show how repeated selective provision of evidence can manipulate learning in the long run.

Our work also contributes to the literature on Bayesian Persuasion started with Kamenica and Gentzkow (2011), where the sender is limited in the amount of information he can provide (see e.g., Bird and Neeman, 2022; Eilat, Eliaz and Mu, 2021; Galperti and Perego, 2024). We add to this body of work by showing that a sender that is subject to a baseline forced-reporting rule can replicate, in the long run, the outcomes of an unconstrained sender and, in certain cases, can even attain the optimal persuasion outcome in equilibrium.

Finally, long-run manipulation of beliefs is also the focus of an extensive literature that studies reputation in dynamic games (see e.g., Kreps and Wilson 1982; Milgrom and Roberts 1982; Fudenberg and Levine 2009; Pei 2020; Ekmekci et al. 2022). In contrast to our paper, in that literature it is typically assumed that it is costly for the agent to take the action that generates the desirable reputation and it is clear which reputation the agent would like to acquire.

References

- Antic, Nemanja, and Archishman Chakraborty.** 2024. “Selected Facts.” CMS-EMS Discussion paper #1608.
- Antic, Nemanja, Archishman Chakraborty, and Rick Harbaugh.** 2024. “Subversive Conversations.” *Journal of Political Economy*, Forthcoming.
- Antler, Yair, Daniel Bird, and Santiago Oliveros.** 2023. “Sequential learning.” *American Economic Journal: Microeconomics*, 15(1): 399–433.
- Arieli, Itai, Yakov Babichenko, Dimitry Shaiderman, and Xianwen Shi.** 2024. “Persuading while Learning.”

- Aumann, Robert J, and Sergiu Hart.** 2003. "Long cheap talk." *Econometrica*, 71(6): 1619–1660.
- Beyer, Anne, Daniel A. Cohen, Thomas Z. Lys, and Beverly R. Walther.** 2010. "The financial reporting environment: Review of the recent literature." *Journal of Accounting and Economics*, 50(2): 296–343.
- Bird, Daniel, and Zvika Neeman.** 2022. "What should a firm know? Protecting consumers' privacy rents." *American Economic Journal: Microeconomics*, 14(4): 257–295.
- Di Tillio, Alfredo, Marco Ottaviani, and Peter Norman Sørensen.** 2021. "Strategic sample selection." *Econometrica*, 89(2): 911–953.
- Dye, Ronald A.** 1985. "Disclosure of nonproprietary information." *Journal of Accounting Research*, 123–145.
- Dye, Ronald A.** 1988. "Earnings management in an overlapping generations model." *Journal of Accounting Research*, 195–235.
- Eilat, Ran, Kfir Eliaz, and Xiaosheng Mu.** 2021. "Bayesian privacy." *Theoretical Economics*, 16(4): 1557–1603.
- Ekmekci, Mehmet, Leandro Gorno, Lucas Maestri, Jian Sun, and Dong Wei.** 2022. "Learning from Manipulable Signals." *American Economic Review*, 112(12): 3995–4040.
- Eliaz, Kfir, Daniel Fershtman, and Alexander Frug.** 2024. "Clerks." Working paper.
- Escobar, Juan F, and Qiaoxi Zhang.** 2021. "Delegating learning." *Theoretical Economics*, 16(2): 571–603.
- Fishman, Michael J, and Kathleen M Hagerty.** 1990. "The optimal amount of discretion to allow in disclosure." *The Quarterly Journal of Economics*, 105(2): 427–444.
- Fudenberg, Drew, and David K Levine.** 2009. "Reputation and equilibrium selection in games with a patient player." In *A Long-Run Collaboration On Long-Run Games*. 123–142. World Scientific.

- Galperti, Simone, and Jacopo Perego.** 2024. "Games with information constraints: seeds and spillovers." *Theoretical Economics*.
- Glazer, Jacob, and Ariel Rubinstein.** 2004. "On optimal rules of persuasion." *Econometrica*, 72(6): 1715–1736.
- Glazer, Jacob, and Ariel Rubinstein.** 2006. "A study in the pragmatics of persuasion: a game theoretical approach." *Theoretical Economics*, 1: 395–410.
- Grossman, Sanford J.** 1981. "The informational role of warranties and private disclosure about product quality." *Journal of Law and Economics*, 24(3): 461–483.
- Grossman, Sanford J, and Oliver D Hart.** 1980. "Disclosure laws and takeover bids." *Journal of Finance*, 35(2): 323–334.
- Healy, Paul M, and James M Wahlen.** 1999. "A review of the earnings management literature and its implications for standard setting." *Accounting Horizons*, 13(4): 365–383.
- Jung, Woon-Oh, and Young K Kwon.** 1988. "Disclosure when the market is unsure of information endowment of managers." *Journal of Accounting Research*, 146–153.
- Kamenica, Emir, and Matthew Gentzkow.** 2011. "Bayesian persuasion." *American Economic Review*, 101(6): 2590–2615.
- Kreps, David M, and Robert Wilson.** 1982. "Reputation and imperfect information." *Journal of Economic Theory*, 27(2): 253–279.
- Krishna, Vijay, and John Morgan.** 2004. "The art of conversation: eliciting information from experts through multi-stage communication." *Journal of Economic Theory*, 117(2): 147–179.
- Milgrom, Paul, and John Roberts.** 1982. "Predation, reputation, and entry deterrence." *Journal of Economic Theory*, 27(2): 280–312.
- Milgrom, Paul R.** 1981. "Good news and bad news: Representation theorems and applications." *Bell Journal of Economics*, 380–391.
- Pei, Harry.** 2020. "Reputation effects under interdependent values." *Econometrica*, 88(5): 2175–2202.

Song Shin, Hyun. 2003. "Disclosures and asset returns." *Econometrica*, 71(1): 105–133.

Verrecchia, Robert E. 1983. "Discretionary disclosure." *Journal of Accounting and Economics*, 5: 179–194.

A Proofs

Proof of Proposition 1. If $\mu_0 \geq v^*$, then by using a strategy profile under which the periodic disclosure is uninformative, all agent-types are accepted with probability one in all periods. That is, all agent-types attain their maximal payoff in the interaction, and so this strategy profile must be an equilibrium.

Next, we consider the case where $\mu_0 \leq v_{-1}^*$. We begin by formally constructing the equilibrium candidate.

Definition of K-Mixing Early-Sanitization Strategy

If the decision-maker's belief that the agent is good is at least v^* , then the agent uses an uninformative disclosure strategy profile. For lower beliefs, we define the equilibrium candidate via a three-dimensional state space $\langle \mu, \tau, k \rangle$, where $\mu \in [0, v^*)$ is the decision-maker's belief that the agent is good, $\tau \in \mathbb{N} \cup \{0\}$ is the duration of the current waiting phase, and $k \in \mathbb{N} \cup \{0\}$ counts the maximal number of waiting phases that can occur in the future (while still outside the acceptance region). The initial state is given by $(\mu_0, 0, K)$.

We now specify the agent's disclosure strategy as a function of the state, and the evolution of the state over time. For each state $\langle \mu, \tau, k \rangle$ and agent's strategy described below, denote by μ'_i the belief derived via Bayes updating following the disclosure of signal i . Furthermore, let $\Xi \equiv (v_{-2}^*, v_{-1}^*)$, where v_{-1}^* is formally defined by

$$v^* = \frac{v_{-1}^*(1 - (1 - \theta_G)^2)}{v_{-1}^*(1 - (1 - \theta_G)^2) + (1 - v_{-1}^*)(1 - (1 - \theta_B)^2)},$$

and v_{-2}^* is defined in analogous manner. The state space can be partitioned into three classes.

1. *Waiting states* are states of the form $\langle \mu, \tau, k \rangle$, where $\tau > 0$. In such states the agent uses an uninformative disclosure strategy profile. Following each waiting period, the state transitions to state $\langle \mu, \tau - 1, k \rangle$.
2. *Sanitization states* are states of the form $\langle \mu, 0, k \rangle$, where $\mu \notin \Xi$. In such states the agent uses within-period sanitization, regardless of his type, and the state transitions to state $\langle \mu'_i, 0, k \rangle$.
3. *Mixing states* are states of the form $\langle \mu, 0, k \rangle$, where $\mu \in \Xi$. In such states, a good agent that observes $\{bg\}$ discloses g with probability $\alpha^*(\mu)$ (defined below), and a bad agent discloses g if possible. Following the disclosure of b , the state transitions to state $\langle \mu'_b, \sigma_b^*(\mu, k), k - 1 \rangle$, whereas following the disclosure of g the state transitions to state $\langle \nu_{-1}^*, \tau^*(\mu, k) + \sigma_g^*(\mu, k), k - 1 \rangle$, where $\sigma_i^*(\cdot, \cdot)$ and $\tau^*(\cdot, \cdot)$ are lotteries with a support of at most two consecutive integers (defined below).

Note that, by construction, since $\mu_0 < \nu_{-1}^*$ the decision-maker's beliefs cannot reach the interval (ν_{-1}^*, ν^*) while $k > 0$.

Definition of the Mixing Probabilities. In this step we define the function

$$\alpha^* : \Xi \rightarrow (0, 1)$$

that specifies the probability that a good agent that observes $\{bg\}$ discloses g in mixing states. Consider a mixing state with belief $\mu \in \Xi$. If a good agent discloses g whenever possible (i.e., $\alpha = 1$), then, according to the definition of Ξ , the decision-maker's belief following the disclosure of g will exceed ν_{-1}^* . Conversely, if a good agent discloses b whenever possible (i.e., $\alpha = 0$), then, under condition KID the decision-maker's beliefs will (weakly) decrease following the disclosure of g . Since the decision-maker's updated belief is continuous in α , by the Intermediate Value Theorem there exists $\alpha^*(\mu) \in (0, 1)$ such that the decision-maker's belief following the disclosure of g is exactly ν_{-1}^* .

Duration of the Waiting Phases. In this step we define the functions

$$\tau^* : \Xi \times \mathbb{N} \rightarrow \tilde{\Delta} \quad ; \quad \sigma_i^* : \Xi \times \mathbb{N} \rightarrow \tilde{\Delta}$$

where $\tilde{\Delta}$ is the space of lotteries with a support of at most two consecutive integers. We construct these functions in K iterative steps.

Step 1: Suppose we are in the last mixing phase with belief $\mu \in \Xi$, i.e., let $k = 1$. We begin by specifying the waiting time $\tau^*(\mu, 1)$ that is designed to make a good agent indifferent between disclosing either signal at mixing state $\langle \mu, 0, 1 \rangle$, under the assumption that $\sigma_i^*(\cdot, 1) = 0$. By Bayes-plausibility the disclosure of b must reduce the decision-maker's belief if a good agent uses the strategy $\alpha^*(\mu)$. If $\tau^*(\mu, 1) = 0$, then a good agent prefers to disclose g whenever possible. Conversely, if $\tau^*(\mu, 1) \rightarrow \infty$, then a good agent prefers to disclose b whenever possible. Let $\bar{\tau}$ denote the maximal integer such that a good agent weakly prefers to disclose g for all waiting phases of duration $\tau(\mu, 1) \leq \bar{\tau}$. If a good agent is indifferent between disclosing b and g for a waiting time of $\bar{\tau}$, set $\tau^*(\mu, 1) = \bar{\tau}$. Otherwise, there exists a lottery with support $\{\bar{\tau}, \bar{\tau} + 1\}$ that makes a good agent indifferent. In this case, set $\tau^*(\mu, 1)$ to be this lottery.

Next, we specify the additional waiting time following the disclosure of g , $\sigma_g^*(\mu, 1)$. This lottery is chosen so that the expected discounted amount of time it takes a good agent that discloses g in a mixing state to reach the acceptance region is the same across all mixing states with $k = 1$. As we show below, this will be useful to incentivize the disclosure of g in sanitization states prior to reaching this mixing state.

Formally, define

$$\hat{\tau}(1) = \inf_{\tilde{\mu} \in \Xi} \{\mathbb{E}(\beta^{\tau^*(\tilde{\mu}, 1)})\},$$

and define $\sigma_g^*(\mu, 1) \in \tilde{\Delta}$ implicitly by

$$\hat{\tau}(1) = \mathbb{E}(\beta^{\tau^*(\mu, 1) + \sigma_g^*(\mu, 1)}).$$

To maintain a good agent's indifference at mixing state $\langle \mu, 0, 1 \rangle$ we must also add a waiting phase following the disclosure of b . Let $\sigma_b^*(\mu, 1) \in \tilde{\Delta}$ be the lottery over waiting times that attains this indifference.¹⁸ This completes step 1.

After $n < K$ steps, we have defined $\tau^*(\mu, k)$ and $\sigma_i^*(\mu, k)$ for all $\mu \in \Xi$ and $k \leq n$. We now consider step $n + 1$. Using $\{\tau^*(\mu, k), \sigma_i^*(\mu, k) | k \leq n, \mu \in \Xi, i \in \{b, g\}\}$, by the same arguments as before there exists a lottery $\tau^*(\mu, n + 1) \in \tilde{\Delta}$ that makes a good agent indifferent between disclosing b or g at mixing state $\langle \mu, 0, n + 1 \rangle$, provided that $\sigma_i^*(\mu, n + 1) = 0$. After defining $\tau^*(\mu, n + 1)$, we define $\sigma_i^*(\mu, n + 1)$

¹⁸The existence of $\sigma_i^*(\mu, 1) \in \tilde{\Delta}$ follows from the same reasons that there exists $\tau^*(\mu, 1) \in \tilde{\Delta}$.

analogously to the way we did in step 1.

Equilibrium

Having defined the equilibrium candidate, we now show that, for every $K \in \mathbb{N}$, the K -mixing early-sanitization strategy is an equilibrium.

In waiting states, the periodic disclosure is meaningless, and so any periodic disclosure strategy is optimal. Similarly, in sanitization phases of the form $\langle \mu, 0, 0 \rangle$ the continuation strategy profile is early sanitization, which we already established to be a continuation equilibrium.

We proceed by backward induction over k .

Step 1: At mixing state $\langle \mu, 0, 1 \rangle$ a good agent is, by the construction of the lotteries $\tau^*(\mu, 1)$ and $\sigma_i^*(\mu, 1)$, indifferent between disclosing either signal. A bad agent, on the other hand, strictly prefers to disclose g . To see this, note that a bad agent is likely to have less g signal realizations than a good agent in the future, and thus his opportunity cost from forgoing the chance to increase the decision-maker's belief is greater than the opportunity cost of a good agent. This implies that the proposed strategy is optimal at mixing states of the form $\langle \mu, 0, 1 \rangle$.

Next we consider sanitization phases of the form $\langle \mu, 0, 1 \rangle$. By the construction of the waiting time lotteries, a good agent's continuation utility is identical at all mixing states of the form $\langle \mu, 0, 1 \rangle$. Moreover, his continuation utility at each such state is positive, whereas his periodic payoff at sanitization phases is zero. It follows that at any sanitization state $\langle \mu, 0, 1 \rangle$, the objective of a good agent is to reach a mixing state as quickly as possible. Since posterior beliefs are increasing in the prior, within-period sanitization minimizes the time – in a first-order stochastic sense – it takes the agent to reach the next mixing state. As a bad agent's cost of forgoing the opportunity to disclose g is greater than a good agent's cost of doing so, within-period sanitization is optimal for both agent types at any sanitization phase with $k = 1$.

Step n : Given $\tau^*(\cdot, \cdot)$ and $\sigma_i^*(\cdot, \cdot)$, for any given $k < n$ the expected time it takes a good agent to get from a mixing state $\langle \cdot, 0, k \rangle$ to the acceptance region is independent of $\mu \in \Xi$. Thus, at any sanitization state $\langle \cdot, 0, k \rangle$, the agent's objective is to reach the next mixing state as quickly as possible. By the definition of $\tau^*(\mu, n)$ and $\sigma_i^*(\mu, n)$, from the perspective of a good agent, at every mixing state of the form $\langle \mu, 0, n \rangle$, disclosing either g or b leads to expected discounted time until the next mixing state

is reached (i.e., a state of the form $\langle \mu, 0, n - 1 \rangle$ where $\mu \in \Xi$), and thus, until a good agent enters the acceptance region. Hence, as before, his objective is to arrive at a mixing state $\langle \mu, 0, n - 1 \rangle$ as soon as possible (the belief $\mu \in \Xi$ at which he enters is not important). Therefore, he will disclose g whenever possible, and by the same single-crossing type of argument we used in step 1, it is in the best interest of a bad agent type to disclose g (whenever possible) during the sanitization states as well.

Maximal Persuasion

Finally, we show that the acceptance rate of a bad agent under these equilibria approaches the KGB as $k \rightarrow \infty$.

Under any K -mixing early-sanitization equilibrium, a good agent is (eventually) accepted with probability one. To see this, note that, by construction, it is optimal for a good agent to disclose g whenever possible. By the Borel-Cantelli Lemma, such an agent will eventually be accepted. This occurs either by disclosing g following a waiting phase or once the continuation strategy reverts to early sanitization (i.e., once $k = 0$).

Moreover, the probability that a good agent will pass through K mixing phases converges to zero as K increases to infinity. To see this, note that with a probability of $\theta_G^4 > 0$, a good agent will have only g signals in two given periods. It follows that with strictly positive probability a good agent will be forced to disclose g both in a mixing state and in the next sanitization state (that occurs after $\tau^* + \sigma_g^*$ periods). Hence, the probability that a good agent will pass through K mixing phases is bounded from above by $(1 - \theta_G^4)^K$; an expression that converges to zero as $K \rightarrow \infty$. By construction, if the agent is accepted in a state with $k > 0$, he enters the acceptance region at belief ν^* . Hence, the distribution of the decision-maker's belief upon entry to the acceptance region converges to a degenerate belief on ν^* as $K \rightarrow \infty$. ■

Proof of Proposition 2.

Part one ($1 \Rightarrow 2$): Let σ_U denote the strategy under which an agent-type with evidence $\{bg\}$ discloses each signal realization with probability $\frac{1}{2}$, and let σ_{BC} denote the belief-contrarian strategy.

Next, define the strategy σ_λ to be the mixture between σ_{BC} with probability λ ,

and σ_U with probability $1 - \lambda$. Finally, let $\psi : [0, 1] \rightarrow \mathbb{R}$ be the function

$$\psi(\lambda) = q_G(\lambda) - q_B(\lambda),$$

where $q_\omega(\lambda)$ is the probability that signal g is disclosed under the strategy σ_λ when the state is ω .

By the premise (part 1 of the lemma) we have that $\psi(1) \leq 0$. On the other hand, under σ_U the disclosure of signal realization g is more likely in state G than in state B . That is, $\psi(0) > 0$.

Note that $q_\omega(\cdot)$ is continuous, and so $\psi(\cdot)$ is also continuous. Hence, by the Intermediate Value Theorem there exists $\lambda^* \in (0, 1]$ such that $\psi(\lambda^*) = 0$. That is, the strategy σ_{λ^*} makes the disclosed signal realization uninformative about the state.

Part two ($2 \Rightarrow 1$) : To establish this part of the lemma, we first establish an auxiliary result. Namely, that for any weak meaning reversal strategy, σ , (other than the belief-contrarian strategy), altering σ so that agent-type $\hat{y}$ with evidence $\{bg\}$ and belief $\mu(\hat{y})$ goes against his belief – discloses $g(b)$ if $\mu(\hat{y}) < (>)\mu_0$ – maintains the weak meaning reversal property.

To see why this auxiliary result is true, fix σ and an agent-type $\hat{y}$ such that: 1) $\hat{y}$ has evidence $\{bg\}$, 2) $\mu(\hat{y}) > \mu_0$ (the opposite case is symmetric), and 3) $\hat{y}$ discloses g with probability $\alpha > 0$. Denote the decision-maker's belief under σ by $P_\sigma(\cdot)$. Let σ' be the strategy that is identical to σ for all $y \neq \hat{y}$, and $\hat{y}$ discloses b with probability one (goes against his belief).

$$\begin{aligned} \mu_0 \geq P_\sigma(G|g) &= \frac{P_\sigma(G \wedge g)}{P_\sigma(g)} \\ &= \frac{\Pr(y \neq \hat{y})P_\sigma(G \wedge g|y \neq \hat{y}) + \Pr(y = \hat{y})P_\sigma(G \wedge g|y = \hat{y})}{\Pr(y \neq \hat{y})P_\sigma(g|y \neq \hat{y}) + \Pr(y = \hat{y})P_\sigma(g|y = \hat{y})} \\ &= \frac{\Pr(y \neq \hat{y})P_\sigma(G \wedge g|y \neq \hat{y}) + \Pr(y = \hat{y})\mu(\hat{y})\alpha}{\Pr(y \neq \hat{y})P_\sigma(g|y \neq \hat{y}) + \Pr(y = \hat{y})\alpha} \\ &> \frac{\Pr(y \neq \hat{y})P_\sigma(G \wedge g|y \neq \hat{y}) + \Pr(y = \hat{y})\mu_0\alpha}{\Pr(y \neq \hat{y})P_\sigma(g|y \neq \hat{y}) + \Pr(y = \hat{y})\alpha}. \end{aligned}$$

The first inequality follows from the assumption that σ reverses the meaning of the signal, and the second inequality follows from the assumption that $\mu(\hat{y}) > \mu_0$. By

simple algebra, the above inequality implies that

$$\mu_0 > \frac{\Pr(y \neq \hat{y})P_\sigma(G \wedge g|y \neq \hat{y})}{\Pr(y \neq \hat{y})P_\sigma(g|y \neq \hat{y})} \equiv P_{\sigma'}(G|g).$$

This establishes the auxiliary result.

Now, assume that a strategy $\hat{\sigma}$ makes the disclosure uninformative. By definition, $\hat{\sigma}$ satisfies the weak meaning reversal property. The auxiliary result established above implies that if we iteratively alter $\hat{\sigma}$ so that every agent-type goes against her belief, the resulting strategy will not only be the belief-contrarian strategy, but will also satisfy the weak meaning reversal property. ■

Proof of Claim 1.

Denote by d the disclosed signal. Under the belief-contrarian strategy

$$\Pr(d = g|\omega) = \Pr(gg|\omega) + \Pr(bg, D|\omega)$$

Weak meaning reversal requires that the likelihood ratio $\frac{\Pr(d=g|G)}{\Pr(d=g|B)}$ be weakly less than one. Plugging the above probability into $\frac{\Pr(d=g|G)}{\Pr(d=g|B)} \leq 1$ and rearranging yields Condition (1).

Note that the left-hand side of Equation (1) is independent of μ_0 . To establish that the right-hand side of this expression is also independent of μ_0 observe that

$$\Pr(bg, D|\omega) = \Pr(bg, \mu(y) \leq \mu_0|\omega) = \Pr(bg|\omega) \Pr(\mu(y) \leq \mu_0|bg, \omega).$$

By definition, $\Pr(bg|\omega)$ is independent of μ_0 . To show that $\Pr(\mu(y) \leq \mu_0|bg, \omega)$ is independent of μ_0 , note that (conditional on evidence $\{bg\}$) a private signal realization of s weakly decreases the agent's beliefs only if $\frac{\Pr(bg, s|G)}{\Pr(bg, s|B)} \leq 1$. Since the evidence and private signal are conditionally independent, this is equivalent to

$$\frac{\Pr(bg|G) \Pr(s|G)}{\Pr(bg|B) \Pr(s|B)} \leq 1 \Leftrightarrow \frac{\Pr(s|G)}{\Pr(s|B)} \leq \frac{\Pr(bg|B)}{\Pr(bg|G)}.$$

Therefore, D is independent of μ_0 . ■

Proof of Proposition 3.

We begin by considering a single period auxiliary setting in which the agent

observes K private realizations of the baseline signal at the beginning of the interaction. Denote by μ_K the agent's (stochastic) belief after observing K private signal realizations. From the law of large numbers it follows that,

$$\lim_{K \rightarrow \infty} \Pr(\mu_K < \mu_0 | G) \rightarrow 0 \text{ and } \lim_{K \rightarrow \infty} \Pr(\mu_K < \mu_0 | B) \rightarrow 1.$$

Denoting by μ_{K+2} the agent's (stochastic) belief after observing the K private signal realizations and the 2 disclosable signal realizations, it follows that

$$\lim_{K \rightarrow \infty} \Pr(gb, \mu_{K+2} < \mu_0 | B) = \lim_{K \rightarrow \infty} \Pr(gb | B) \Pr(\mu_{K+2} < \mu_0 | gb, B) = 2\theta_B(1 - \theta_B),$$

and that

$$\lim_{K \rightarrow \infty} \Pr(gb, \mu_{K+2} < \mu_0 | G) = \lim_{K \rightarrow \infty} \Pr(gb | G) \Pr(\mu_{K+2} < \mu_0 | gb, G) = 0.$$

The LHS of (1) is $\theta_G^2 - \theta_B^2$ regardless of K , and the above calculations show that the RHS of (1) converges to $2\theta_B(1 - \theta_B)$ as $K \rightarrow \infty$. Thus, Condition (S-KID) implies that Equation (1) is satisfied for large enough K . In combination with Claim 1, this implies that there exists K^* such that there exists an uninformative disclosure strategy for all $K > K^*$.

The above implies that in the dynamic model, where the agent is initially uninformed, he can stop the decision-maker's learning after period $T^* = K^*$. In particular, assume that for T^* periods the agent chooses at random which of the two signal realizations to disclose. Conditional on the state, the T^* concealed draws are independent of the disclosed history (under random disclosure), so at the beginning of period $T^* + 1$ the agent's private information is equivalent to K^* additional i.i.d. baseline draws relative to the decision-maker. Since Condition (1) (hence the existence of an uninformative disclosure profile) is prior-independent, an uninformative disclosure strategy exists at period $T^* + 1$ for any history and can be repeated forever (ignoring future information if needed).

Next, we show that if (S-KID) does not hold, then the disclosed signal realization is informative in every period. To see this, note that at every period the sequences of periodic signal realizations $(gg), (gg), \dots, (gg)$ and $(bb), (bb), \dots, (bb)$ occur with positive probability in both states of nature. Thus, at every period there is a positive probability that the agent updated his belief in the wrong direction (i.e., increases

his belief relative to the prior when the state is B and vice versa). This, in turn, implies that regardless of the agent's strategy, $\Pr(gb, \mu(y) < \mu_0|B) - \Pr(gb, \mu(y) < \mu_0|G) < 2\theta_B(1 - \theta_B)$. Hence, Condition (1) does not hold, and by Proposition 2 an uninformative disclosure strategy does not exist. ■

Proof of Proposition 4. Let $\hat{f}_T \in \Delta(R)$ be the empirical distribution of reports up to time T . For each report $r \in R$ define

$$m_t(r) \equiv \mathbb{E}(\mathbb{1}(\hat{\sigma}_t = r)|h_{t-1}, \omega).$$

Note that $\mathbb{1}(\hat{\sigma}_t = r) - m_t(r)$ is a bounded martingale difference sequence, so by the strong law for martingale differences

$$\lim_{T \rightarrow \infty} \left(\max_{r \in R} |\hat{f}_T(r) - \frac{1}{T} \sum_{t=1}^T m_t(r)| \right) \rightarrow 0 \text{ a.s.}$$

For each t , $m_t(\cdot) = f(\cdot|\omega, \sigma_t) \in \mathcal{F}(\omega)$. Since $\mathcal{F}(\omega)$ is convex and closed,

$$1/T \sum_{t=1}^T m_t(\cdot) \in \mathcal{F}(\omega),$$

and therefore the distance between $\hat{f}_T$ and $\mathcal{F}(\omega)$ converges almost surely to zero as $T \rightarrow \infty$.

Because the sets $\{\mathcal{F}(\omega)\}_{\omega \in \Omega}$ are pairwise disjoint and compact subsets of $\Delta(R)$, they are separated by a positive distance; hence for large T , $\hat{f}_T$ is (with probability tending to one) close only to $\mathcal{F}(\omega')$, the true-state set. This implies the likelihood of the observed frequency under any $\omega \neq \omega'$ vanishes, so that the DM's posterior converges to $\delta_{\omega'}$. ■

Proof of Proposition 5.

Fix a strictly manipulable set S and a finite support belief distribution $\{p_k, \nu_k\}_{k=1}^K$ that is Bayes-plausible conditional on S .

We begin by constructing a specific reporting strategy. Fix an arbitrary state-contingent reporting strategy σ_0 . For states $\omega \in \Omega \setminus S$ the agent uses σ_0 in every period. To construct the rest of the agent's strategy, we need the following notation.

For each $\omega \in \Omega \setminus S$ and $\epsilon > 0$, let $b_\epsilon(\omega) \subset \Delta(R)$ denote the open ϵ ball around

the long-run reporting frequency $f(\cdot|\omega, \sigma_0)$. Let

$$B_\epsilon = \bigcup_{\omega \in \Omega \setminus S} b_\epsilon(\omega).$$

Fix ϵ for which the set

$$H = \bigcap_{\omega \in S} \mathcal{F}(\omega) \setminus B_\epsilon$$

contains a line of strictly positive length. Such ϵ exists since S is strictly manipulable.

To construct the agent's strategy in S , select K distinct reporting distributions $\{r_k\}_{k=1}^K$ that belong to H . For each such distribution, denote the corresponding state-contingent reporting strategy in state $\omega \in S$ by $\sigma_k(\omega)$. At period $t = 0$, conditional on $\omega \in S$, the agent draws an index $k \in \{1, \dots, K\}$ with probabilities $\frac{p_k v_k(\omega)}{\mu_0^S(\omega)}$. This randomization is well defined as Bayes-plausibility conditional on S implies that

$$\sum_{k=1}^K q_k(\omega) = \frac{\sum_{k=1}^K p_k v_k(\omega)}{\mu_0^S(\omega)} = 1.$$

The agent then uses the strategy $\sigma_k(\omega)$ in all periods.

The probability that the agent draws k is

$$\Pr(k) = \sum_{\omega \in S} \mu_0^S(\omega) \frac{p_k v_k(\omega)}{\mu_0^S(\omega)} = \sum_{\omega \in S} p_k v_k(\omega) = p_k.$$

The Law of Large Numbers implies that the distribution of the long-term reporting frequencies will converge to the report distribution conditional on the realized state and the chosen reporting strategy. By construction, for all $\omega \notin S$ the distance between H and $f(\cdot|\omega, \sigma_0)$ is at least ϵ . Hence, by the LLN, the decision-maker will be able to infer whether or not the state belongs to S . Moreover, after observing a long-term reporting frequency of $r_k \in H$ the decision-maker's belief of the state being ω is

$$\frac{\mu_0^S(\omega) \frac{p_k v_k(\omega)}{\mu_0^S(\omega)}}{\sum_{\omega' \in S} (\mu_0^S(\omega') \frac{p_k v_k(\omega')}{\mu_0^S(\omega')})} = \frac{p_k v_k(\omega)}{\sum_{\omega' \in S} (p_k v_k(\omega'))} = \frac{v_k(\omega)}{\sum_{\omega' \in S} v_k(\omega')} = v_k(\omega).$$

■

B Reputation management with $\mu \in (\nu_{-1}^*, \nu^*)$

In this appendix we revisit the reputation management story. We begin by showing that for our specification of the decision-maker's behavior, while the manipulation in equilibrium is significant (e.g., via early sanitization strategy), the KGB cannot be approximated when the decision-maker's prior belief is in (ν_{-1}^*, ν^*) . We conclude with a remark on a possible modification to a history-dependent decision-maker's strategy (interpreting preferences at the threshold ν^* as indifference), under which KGB could be attained for prior beliefs in (ν_{-1}^*, ν^*) as well. While this approach is even more permissive, it relies on a rather complex cooperation on the decision-maker's part, which we find less natural or robust than a choice rule which depends only on the current belief.

Proposition 6. *If $\mu_0 \in (\nu_{-1}^*, \nu^*)$, then the KGB cannot be approximated in equilibrium.*

Proof. Any equilibrium disclosure strategy σ defines a probability space over binary sequences of disclosed evidence $\{b, g\}$, and a corresponding martingale of end-of-period beliefs $\mu_\sigma(h_t)$ following the disclosure in the first t periods. Let $\hat{H}(\sigma)$ be the set of (infinite) histories along which beliefs never fall below $\frac{\mu_0 + \nu_{-1}^*}{2}$ under σ . Since beliefs are a martingale with a support that is bounded in $[0, 1]$, Doob's Martingale Inequality implies that

$$\Pr(\hat{H}(\sigma)) \geq \frac{\mu_0 - \nu_{-1}^*}{2 - \mu_0 - \nu_{-1}^*} \equiv M_1 > 0.$$

Fix an equilibrium disclosure strategy. Conditioning on $\hat{H}(\sigma)$, the following three events form a partition; hence at least one occurs with conditional probability at least $1/3$ (and therefore with unconditional probability at least $M_1/3$).

i): A history in $\hat{H}$ never enters the AR (acceptance region: $\mu(\cdot) \geq \nu^*$). In this case, the outcome is bounded away from KGB since a good agent is never accepted with a probability of at least $\frac{M_1}{3} \frac{\mu_0 + \nu_{-1}^*}{2}$.

ii): Under a history in $\hat{H}$ the first entry to AR occurs at a period where a good agent uses within-period sanitization. In this case, the outcome is bounded away from the KGB since the belief following entry to AR is at least $\bar{\nu} > \nu^*$ (a belief that results following within-period sanitization that begins at belief $\frac{\mu_0 + \nu_{-1}^*}{2}$). Note that, at such histories, the report g is made with a probability of at least θ_B^2 , so that the decision-maker's beliefs upon entry to AR will be at least $\bar{\nu}$ with a probability of at

least $\theta_B^2 \frac{M_1}{3}$. Following such beliefs at entry, Doob's Martingale inequality implies that the decision-maker's long-run belief will be above $\frac{\bar{v} + v^*}{2}$ with probability of at least $M_2 \equiv \frac{\bar{v} - v^*}{2 - \bar{v} - v^*}$. Thus, the long-run beliefs are bounded away from the KGB.

iii): Under a history in $\hat{H}$ the first entry to AR occurs when a good agent that observes $\{b, g\}$ discloses b with positive probability. Denote all the prefixes at which such entrance to AR occurs by $\mathcal{H}_{b \geq g}^{AR}$.

For any $V < \frac{1}{1-\beta}$, let $\bar{T}(V)$ denote the minimal integer such that

$$\sum_{t=0}^{\bar{T}(V)} \beta^t \geq V.$$

Moreover, let $\kappa \equiv \min\{(1 - \theta_G)^2, \theta_B^2\}$. Intuitively, κ is a uniform lower bound on the probability that a given signal is disclosed in a specific period. Next, we derive the following lemma about such histories (proof below).

Lemma B.1. *The continuation payoff of both types at any history at $\mathcal{H}_{b \geq g}^{AR}$ is below $\frac{1}{1-\beta} - \Delta^*$, where $\Delta^* \equiv (\kappa\beta)^{\bar{T}(\frac{\beta}{1-\beta})}$.*

For each $h_t \in \mathcal{H}_{b \geq g}^{AR}$, we can define a sequence $u(h_t) = (u_j(h_t))_{j=1}^\infty$ such that the decision-maker's belief is non-decreasing along the history $(h_t, u(h_t))$. Such a sequence exists since the belief process is a martingale. We now partition the set

$$\mathcal{H}_{b \geq g}^{AR} = I_1 \cup I_2 \cup \dots \cup I_{T^+},$$

where $T^+ = \bar{T}(\frac{1}{1-\beta} - \Delta^*)$, as follows.

For each $h_t \in \mathcal{H}_{b \geq g}^{AR}$, if in period $t + 1$, following h_t , i) both agent-types disclose evidence according to a pure strategy and ii) disclosure is informative, let $h_t \in I_1$. Otherwise, if in period $t + 2$, i) both agent-types use a pure strategy at $(h_t, u_1(h_t))$ and ii) disclosure is informative, let $h_t \in I_2$. Continuing in this manner, $h_t \in I_{j+1}$ if at every period along $(h_t, (u_i(h_t))_{i=1}^{j-1})$ either disclosure was noninformative or at least one agent-type uses a mixed strategy, whereas at $(h_t, (u_i(h_t))_{i=1}^j)$ disclosure is informative and both agent types use a pure strategy.

By Lemma B.1 and the definition of T^+ ,

$$h_t \in I_1 \cup I_2 \cup \dots \cup I_{T^+}.$$

Note that there exists $S \in \{1, \dots, T^+\}$ for which the measure of I_S is at least $\frac{M_1 \kappa^{T^+}}{3T^+} > 0$.

If i) the decision-maker's prior belief is μ , ii) both agent types use a pure strategy, and iii) disclosure is informative, then the support of the belief in the following period is disjoint from $[\mu - \epsilon(\mu), \mu + \epsilon(\mu)]$ for some $\epsilon(\mu) > 0$. Furthermore, note that $\epsilon(\mu)$ is continuous in μ . Hence, there exists $\epsilon^* > 0$ such that $\epsilon^* \leq \epsilon(\mu)$ for every $\mu \in [\nu^*, \frac{\nu^*+1}{2}]$. Note that if $\mu \geq \nu^*$, then the decision-maker's belief in the following period is at least $\min\{\nu^* + \epsilon^*, \frac{\nu^*+1}{2}\}$ with a probability of at least κ . Intuitively, exactly one of the signals will lead to an increase in the decision-maker's belief, and κ is a lower bound on the probability that the agent must disclose that signal. It follows that for any $h_t \in I_S$ the decision-maker's belief after the period $t + S$ disclosure is at least $\min\{\nu^* + \epsilon^*, \frac{\nu^*+1}{2}\}$ with a probability of at least κ . Hence, the outcome is bounded away from the KGB outcome. ■

Proof of Lemma B.1. Fix a history in $h_t \in \mathcal{H}_{b \geq g}^{AR}$. Note that at h_t entry to AR occurs following exactly one of the disclosed signal realizations, either b or g .

First, consider the case where entry to AR occurs via the disclosure of g . At period t at a prefix of h_t , a good agent that has evidence $\{b, g\}$ weakly prefers disclosing b . As disclosing b will lead to a payoff of zero in period t , this preference implies that the continuation payoff from disclosing g , from the perspective of that agent, is at most $\sum_{i=1}^{\infty} \beta^i \cdot 1 = \frac{\beta}{1-\beta}$.

This means that, upon entry to AR, from the perspective of a good agent, there is a non-zero vector $(r_1, r_2, \dots)$, where r_j represents the probability of leaving AR (for the first time) at period $t + j$, from the perspective of period t . Moreover, since this continuation payoff is at most $\frac{\beta}{1-\beta}$, the first $\bar{T}(\frac{\beta}{1-\beta})$ components of this vector cannot all equal zero. Note that any nonzero component from the perspective of a good agent also corresponds to a non-zero component from the perspective of a bad agent, and moreover, each such component is bounded from below by $\kappa^{\bar{T}(\frac{\beta}{1-\beta})}$ for both agent types.

Next, consider the case where entry to AR occurs via the disclosure of b . In this case, at period t (the prefix of h_t), a bad agent that has evidence $\{b, g\}$ must disclose g with positive probability. Otherwise, disclosing b leads to a reduction of the decision-maker's belief, contradicting the assumption that AR is entered via the disclosure of b . Thus, reporting g (and staying outside of AR in period t) is optimal for a bad agent, and so his continuation payoff is at most $\frac{\beta}{1-\beta}$. An analogous argument

to the one used in the previous case, establishes the lemma. ■

Remark: To approximate the KGB for priors in (ν_{-1}^*, ν^*) (when condition KID holds) one could modify the decision-maker's behavior at the threshold belief ν^* in a history dependent manner (this can be justified by interpreting the decision-maker's choices as arising from indifference between acceptance and rejection at the threshold). Specifically, when at the threshold belief ν^* , the decision-maker can postpone acceptance for some time (rather than accept the agent immediately). Under this alternative specification, we could use a construction analogous to the one used in Proposition 1 in which the building block of the 1-mixing early sanitization strategy is modified as follows. The type G agent uses a mixed strategy when the decision-maker's belief is an element of (ν_{-1}^*, ν^*) (rather than (ν_{-2}^*, ν_{-1}^*)), and the mixing probability is such that the disclosure of a signal realization of g leads to an updated belief of ν^* (rather than ν_{-1}^*). Note that in order to incentivize such mixing the duration of postponement following the disclosure of g must depend on the continuation of the game following the disclosure of b .

C Weak manipulability

When $|\cap_{\omega \in S} \mathcal{F}(\omega)| = 1$ the construction used in the proof of Proposition 5 cannot be applied. In such cases the overlap is too thin to directly support the richness of frequencies needed both to separate from states outside S and to mix between distinct frequencies within S . Nevertheless, as we illustrate below, the scope of manipulation can still be large. Achieving this, however, typically requires a more complex reporting strategy than the one used in the proof of Proposition 5.

Consider the selective forced disclosure model, and assume that $|\mathcal{F}(G) \cap \mathcal{F}(B)| = 1$, i.e., that $\theta_G^2 = 2\theta_B - \theta_B^2$. Denote the state contingent strategy that induces the same report distribution in both states by σ_U . Moreover, denote the strategy of disclosing either signal at random by σ_I . Note that if the agent uses σ_I the disclosed signal is informative about the state. For $\gamma \in [0, 1]$, denote by $\sigma_M(\gamma)$ the strategy under which the agent uses strategy σ_I with probability γ and σ_U with probability $1 - \gamma$.

Fix a target long-term belief distribution with finite support that satisfies Bayes-plausibility. That is, select $F^* = \{p_k, \nu_k\}_{k=1}^K$ such that $0 \leq \nu_1 < \nu_2 < \dots < \nu_K \leq 1$,

$\sum_k p_k = 1$, and $\sum_k p_k v_k = \mu_0$. To simplify exposition, assume that $\mu_0 \neq v_k$.¹⁹ Define k_0 by $v_{k_0} < \mu_0 < v_{k_0+1}$, and let G_t denote the distribution of beliefs at the beginning of period t .

Targeting: For a belief distribution G_t and two target beliefs $x < y$ such that $\text{supp}(G) \subseteq (x, y)$, we say that the agent targets $\langle x, y \rangle$ if, regardless of the disclosure history, the agent uses $\sigma(\gamma(x, y, G_t))$, where $\gamma(x, y, G_t)$ is the maximal probability such that $\text{supp}(G_{t+1}) \subseteq [x, y]$. Such a strategy exists since for $\gamma = 0$ there is no updating, the updated beliefs are continuous in γ , and the maximal (minimal) element of G_{t+1} is strictly increasing (decreasing) in γ .

The Algorithm: Now we are ready to construct an algorithm that enables the agent to generate beliefs $\{G_t\}_{t=1}^\infty$ that converge (in distribution) to F^* as $t \rightarrow \infty$.

Define the initial targets as $x = v_{k_0}, y = v_{k_0+1}$, the initial capacities as $\kappa_x = p_{k_0}, \kappa_y = p_{k_0+1}$, and let G_1 denote the degenerate belief distribution that assigns probability one to the prior. Let $t = 1$, and apply the following procedure.

Assume that from G_t the agent targets beliefs $\langle x, y \rangle$, and consider two cases: *Case a: no overfilling:* We say that targeting $\langle x, y \rangle$ from beliefs G_t does not lead to overfilling if the mass on $x(y)$ under G_{t+1} is at most $\kappa_x(\kappa_y)$. In this case, the agent uses this targeting strategy in period t . For any history that leads to one of the targets, the belief is frozen. That is, in any continuation of such a history the agent uses σ_U . Furthermore, if the mass on a $x(y)$ is exactly $\kappa_x(\kappa_y)$, then reset the target and capacity to $x = v_{k_0-1}, \kappa_x = p_{k_0-1}(y = v_{k_0+2}, \kappa_y = p_{k_0+2})$. For all non-frozen histories in G_{t+1} , recalculate that targeting strategy for $\langle x, y \rangle$, and repeat.

Case b: Overfilling We focus on the case where the mass of beliefs on y in G_{t+1} exceeds κ_y . The cases where there is overfilling of x or overfilling of both targets are analogous. Note that under the targeting strategy the disclosure of signal realization g leads to an increase of the updated beliefs, and that beliefs are updated to y only from the maximal non-frozen belief in G_t , denote this belief by $\bar{g}_t$.

Let q denote the probability of disclosing g under the targeting strategy $\gamma(x, y, G_t)$ from belief $\bar{g}_t$. Since σ_I is interior in both states of nature, the reporting strategy can be changed so that beliefs following the disclosure of g remain y but the probability of disclosing g is any element of $B_\epsilon(q)$, for some $\epsilon > 0$. Moreover, there exists $\delta > 0$ such that for any belief in $B_{\bar{g}_t}(\delta)$ targeting y is possible.

For beliefs in $B_{\bar{g}_t}(\delta)$, the set of probabilities with which signal g can be disclosed

¹⁹It can be shown that this is WLOG.

so that the updated belief following the disclosure of g is exactly y is continuous in the initial belief. Hence, if the agent first uses the strategy $\sigma(\gamma)$ N times, where γ is sufficiently small and N is sufficiently large, and then, from each resulting belief, either targets $\langle x, y \rangle$ or temporarily freezes the belief, he will be able to induce belief y with a probability of exactly κ_y in period $t + N + 1$. Assume that the agent does so, and in period $t + N + 2$ freezes all histories that reach belief y and unfreezes the beliefs that were temporarily frozen in these N periods. Finally, reset the target and capacity to $y = v_{k_0+2}, \kappa_y = p_{k_0+1}$, and for all unfrozen beliefs in G_{t+N+2} recalculate that targeting strategy, and repeat.

Since the updated beliefs converge in distribution to beliefs with support $\{0, 1\}$ if $\gamma = 1$, this process will, almost surely, fill all beliefs apart from v_1 and v_K in finite time. Note that by Bayes-plausibility if all but one of the targets has been achieved, the last target will have been achieved as well. Hence, this algorithm generates a sequence of belief distributions $\{G_t\}_{t=1}^{\infty}$ that converge to the target F^* .