

October 2025

“Mutual Reputation and Trust in a Repeated Sender–Receiver Game”

Georgy Lukyanov

Mutual Reputation and Trust in a Repeated Sender–Receiver Game*

Georgy Lukyanov[†]

Abstract

We study a repeated sender–receiver game where inspections are public but the sender’s action is hidden unless inspected. A detected deception ends the relationship or triggers a finite punishment. We show the public state is low-dimensional and prove existence of a stationary equilibrium with cutoff inspection and monotone deception. The sender’s mixing pins down a closed-form total inspection probability at the cutoff, and a finite punishment phase implements the same cutoffs as termination. We extend to noisy checks, silent audits, and rare public alarms, preserving the Markov structure and continuity as transparency vanishes or becomes full. The model yields testable implications for auditing, certification, and platform governance: tapering inspections with reputation, bunching of terminations after inspection spurts, and sharper cutoffs as temptation rises relative to costs.

Keywords: bilateral reputation; trust; costly verification; auditing; private monitoring; repeated games.

JEL: C73; D82; D83.

*I thank Johannes Hörner and Allen Vong for helpful comments. All errors are my own.

[†]Affiliation: *Toulouse School of Economics*. Email: georgy.lukyanov@tse-fr.eu.

1 Introduction

Tax authorities, financial supervisors, and certification bodies routinely publicize their inspection policies—how often they audit, when they intensify scrutiny, and when they stand down—while the underlying conduct of firms remains unobserved unless an audit uncovers a violation. Across these environments one repeatedly sees the same patterns: inspection intensity tapers after periods of compliant findings, surges of auditing are followed by clusters of detected violations, and regimes that replace “one-strike” termination with finite probation achieve similar deterrence with different administrative costs. This paper provides a microfoundation for these facts by analyzing a repeated sender–receiver game with public inspections and private actions. We establish existence of a stationary equilibrium in a low-dimensional belief state (honesty and vigilance), show that optimal policies take cutoff form (the receiver checks when honesty beliefs are low enough; the sender deceives only when vigilance beliefs are low enough), deliver locally unique thresholds with closed-form expressions in a benchmark, and prove that finite, publicly observed punishments implement the same cutoffs as immediate termination. The comparative statics are transparent: higher temptation expands the checking region, higher inspection cost contracts it, greater patience reduces the inspection intensity needed to sustain honesty; inspection intensity tapers as honesty accumulates, and terminations bunch after audit spurts.

We study a repeated sender–receiver game in which the receiver’s inspection decision is publicly observed while the sender’s action is private unless checked. A detected deception ends the relationship in our benchmark, and—importantly for applications—can instead trigger a finite, publicly observed punishment phase. The public state is low-dimensional: two beliefs summarize history, the receiver’s belief about the sender’s honesty and the sender’s belief about the receiver’s vigilance. Despite the private-action friction, the public observability of inspections makes the equilibrium Markov in these beliefs and supports sharp characterizations.

We prove existence of a stationary Perfect Bayesian equilibrium (Lemma 7.1) and show that equilibrium policies take a cutoff form: the receiver checks when honesty beliefs are low enough and the sender deceives only when vigilance beliefs are low enough. Monotone best responses deliver these thresholds (Lemma 5.2); the cutoffs are pinned down by two indifference equations that equate current costs or benefits with discounted shifts in continuation values (equations (5.1)–(5.2)).

In a simple benchmark that captures short deterrence “windows,” we obtain closed-form expressions for the total inspection rate required to deter and for the receiver’s honesty cutoff (Lemma 5.3). We then establish that finite, visible punishments implement the same cutoffs as immediate termination (Proposition 6.3), showing that the mechanism sustaining discipline is reputational rather than dependent on absorption.

Finally, we derive comparative statics and hazard results: greater temptation expands the checking region while higher inspection costs contract it; greater patience lowers the inspection intensity needed to sustain honesty (Proposition 8.1); inspection intensity tapers as honesty reputation accumulates, and terminations bunch after inspection spurts (Corollary 8.2). These

predictions map directly to settings like regulatory audits, certification, platform moderation, and insurance fraud detection.

Technically, we make four additions that tighten the baseline analysis: we prove value monotonicity and single-crossing in our exact environment (Proposition 5.1), give a self-contained existence and continuity result for stationary PBE (Theorem 8.3), establish local uniqueness of the joint mixing cutoffs via a generalized Jacobian (Lemma 7.2 and Corollary 7.3), and provide robustness for the closed-form inspection equalization through explicit bounds (Proposition 5.6). We also integrate transparency and welfare: a short synthesis shows how noisy verification and silent audits tighten discipline (Proposition 10.1), and a hazard-based welfare decomposition clarifies policy trade-offs (Theorem 9.1).

2 Related Literature

This paper connects classical reputation with incomplete information (Kreps and Wilson, 1982; Milgrom and Roberts, 1982) to repeated games with imperfect (here, one-sided public) monitoring. Our sender–receiver environment features two long-lived players who jointly shape beliefs: the receiver builds a reputation for vigilance while the sender privately chooses whether to deceive. Methodologically, we marry cutoff characterizations and belief-state dynamics with the self-generation approach from repeated games (Abreu et al., 1990).

Relative to the reputation literature with a patient long-run player facing short-run opponents (Fudenberg and Levine, 1989) and more recent work on reputations for honesty (Fudenberg et al., 2022), we study *mutual* (two-sided) reputational incentives under asymmetric observability and endogenous termination. Compared to repeated games with almost-public or private monitoring (Mailath and Morris, 2002; Green and Porter, 1984), our receiver’s public checking creates a tractable public signal while the sender’s action remains privately observed, yielding transparent stationary mixing and cutoff cut-loci.

Our setting also relates to costly verification and auditing (Townsend, 1979; Mookherjee and Png, 1989; Kofman and Lawarrée, 1993), and to dynamic certification/monitor reputation. The model provides new, testable predictions for vigilance cutoffs, experimentation, and termination statistics with applications to platform trust and verification markets, compliance/auditing, and expert oversight.

The paper complements work on reputational cheap talk and persuasion with fact-checking and visibility frictions. In particular, our analysis speaks to recent models of public persuasion with endogenous verification (Lukyanov and Safaryan, 2025) and reputational signaling (Lukyanov, 2023), while focusing on two-sided long-run interaction rather than one-sided market-discipline or certification.

Our contribution connects two strands and departs from both. In the costly verification tradition, monitoring shapes incentives but reputational capital resides with the agent being monitored; classical analyses (e.g., Mookherjee and Png (1992, 1994)) focus on optimal verification without a public reputation for the monitor. In dynamic certification and monitor–reputation models (e.g., Marinovic et al. (2018); Marinovic and Szydlowski (2023)), certifiers build credibility through disclosure, but the certified party’s action is typically

publicly evaluable ex post.

We instead place *both* reputations in play under asymmetric observability: the receiver’s inspections are public (so she accrues vigilance reputation) while the sender’s action is private unless inspected (so he accrues honesty reputation). This asymmetry yields a one-dimensional public state with endogenous stopping, generating the cutoff structure, hazards, and transparency results that differ from standard costly–verification and certification benchmarks.

The paper proceeds as follows. Section 3 sets out the environment, timing, information structure, and the termination and punishment variants. Section 4 defines stationary equilibrium in the public belief state and derives the indifference conditions. Section 5 establishes monotonicity and the cutoff characterization and presents the closed-form benchmark. Section 6 proves the equivalence between termination and finite punishment. Section 7 gives existence and a local uniqueness result. Section 8 develops comparative statics and hazard/welfare implications, and Section 9 discusses applications and testable predictions. Section 10 studies robustness—silent audits, noisy checks, finite horizons, and a continuous-time sketch—and discusses the fully private-monitoring benchmark with minimal public signals. Section 11 concludes.

3 Model

3.1 Players, actions, and payoffs

There is a sender (player 1; she) and a receiver (player 2; he). Time is discrete, $t = 1, 2, \dots$, and both discount future with factor $\delta \in (0, 1)$. Each period the sender chooses privately $a_t^s \in \{\text{truth}, \text{deceive}\}$ and the receiver chooses publicly $a_t^r \in \{\text{trust}, \text{check}\}$. Let $B > 0$ denote the sender’s one-shot benefit from deceiving a trusting receiver, and let $C \in (0, B)$ be the receiver’s cost of checking.

If the receiver trusts, her realized payoff is B when the sender is truthful and 0 when the sender deceives; the sender’s is 0 when truthful and B when deceiving. If the receiver checks, she pays C and perfectly learns the sender’s current-period action (the audit is errorless). Upon a detected deception, either (i) the relationship terminates (benchmark) or (ii) the game continues in a *punishment phase* (extension; see below). If the receiver checks a truthful sender, no termination occurs and payoffs are $(0, B - C)$.

3.2 Types and priors

The sender may be a *committed honest type* who always plays **truth**, with prior probability $\lambda_0 \in (0, 1)$, or a *strategic type* who chooses actions optimally. The receiver may be a *committed vigilant type* who always plays **check**, with prior probability $\mu_0 \in (0, 1)$, or a *strategic type*. Types are independent across players, drawn once at $t = 0$, and commonly known in distribution.

3.3 Information and timing

At the start of period t the public state consists of the publicly observed history of receiver actions and whether the relationship is active. The sender privately chooses a_t^s ; then the receiver publicly chooses a_t^r . If $a_t^r = \text{trust}$, no further public signal is realized in that period and the sender's action remains private. If $a_t^r = \text{check}$, the receiver perfectly observes a_t^s and—if $a_t^s = \text{deceive}$ —deception is detected (triggering termination or punishment). Stage payoffs are realized and (except when detection/termination occurs) not publicly revealed.

3.4 Histories and strategies

Let $h_t^p = (a_1^r, \dots, a_{t-1}^r, \text{active at } t)$ denote the public history at t . The sender's private history h_t^s augments h_t^p with her past actions and realized payoffs; the receiver's private history h_t^r augments h_t^p with any audit findings from past checks. A (behavioral) strategy for the strategic sender is $\sigma : \mathcal{H}^s \rightarrow \Delta(\{\text{truth}, \text{deceive}\})$, and for the strategic receiver $\rho : \mathcal{H}^p \rightarrow \Delta(\{\text{trust}, \text{check}\})$. We will focus on stationary belief-based (Markov) strategies in which σ and ρ depend on a low-dimensional belief state defined below.

3.5 Beliefs and reputation states

Let $\lambda_t \in [0, 1]$ denote the receiver's (public) belief at the start of period t that the sender is the committed honest type, and let $\mu_t \in [0, 1]$ denote the sender's (public) belief that the receiver is the committed vigilant type. Because the receiver's action is public, $\mu_{t+1} = 0$ after any realized $a_t^r = \text{trust}$; after $a_t^r = \text{check}$, Bayes' rule (given the receiver's equilibrium mixing) updates μ upward or downward as appropriate. The receiver updates λ only following checks: if a check reveals **truth**, λ rises; if it reveals **deceive**, λ collapses to 0 and either the relationship ends (benchmark) or a punishment phase starts (extension). When the receiver trusts, λ remains at its prior value because the sender's action is unobserved. Along the equilibrium path with stationary mixing, the public belief state $x_t := (\lambda_t, \mu_t)$ evolves as a time-homogeneous Bayesian recursion driven by the observed public action a_t^r and, on check histories, the audit outcome.

3.6 Benchmark: termination upon detected deception

In the benchmark environment, the first period t with $(a_t^s, a_t^r) = (\text{deceive}, \text{check})$ triggers termination immediately after stage payoffs are realized. Let $V^s(x)$ and $V^r(x)$ denote the strategic sender's and receiver's continuation values at public belief state $x = (\lambda, \mu)$. Equilibrium strategies $\sigma(x)$ and $\rho(x)$ will be characterized by indifference conditions that pin down *cutoffs/mixing* in λ and μ .

Table 1: Notation

Symbol	Meaning
$x = (\lambda, \mu)$	Public belief state.
$\sigma(x), \rho(x)$	Sender’s deception probability; receiver’s inspection probability.
B, C, δ	Benefit from deception; inspection cost; discount factor.
truth, deceive	Sender actions (private unless checked).
trust, check	Receiver actions (public).
$p^{\text{check}}(x)$	Probability a check occurs.
$p^{\text{T check}}(x)$	Probability that a check reveals truth .
$\lambda^+(x)$	Honesty posterior after a (truthful) check.
$\mu^+(x)$	Vigilance posterior after a check.
$V_s(x), V_r(x)$	Continuation values (normal phase).
$V_s^{\text{pun}}(x), V_r^{\text{pun}}(x)$	Continuation values upon entering punishment (Section 6).
$p^{\text{check}*}, p^{\text{T check}*}$	Starred probabilities at the mixing cutoffs.
$h(x)$	Hazard of termination (Eq. (8.2)).

3.7 Extension: reputational punishment instead of termination

Upon a detected deception, the game transitions to a publicly observed punishment phase of finite length $T \in \mathbb{N}$ (or an absorbing “bad” state) in which continuation values are reduced via prescribed inspection and trust policies and/or transfers (e.g., automatic checks, exclusion from trade). After punishment, the game returns to the normal phase. We will show (under parameter restrictions) an *outcome-equivalence*: there exists a finite T and public punishment policy that implement the same cutoffs as in the termination benchmark.

3.8 Objective

We seek stationary PBE with the following structure: (i) the receiver’s inspection probability $\rho(\lambda, \mu)$ is monotone in λ ; (ii) the sender’s deception probability $\sigma(\lambda, \mu)$ is 0 above a vigilance cutoff and positive below; and (iii) beliefs (λ_t, μ_t) follow a one-step Bayesian recursion induced by (σ, ρ) . We then derive comparative statics in (B, C, δ) and characterize the hazard of termination (or punishment) and welfare.

The public state is $x = (\lambda, \mu) \in [0, 1]^2$, where λ is the receiver’s belief that the sender is the committed honest type and μ is the sender’s belief that the receiver is the committed vigilant type. The sender’s strategic deception probability is $\sigma(x) \in [0, 1]$; the receiver’s strategic inspection probability is $\rho(x) \in [0, 1]$. The sender’s one-period gain from deception is $B > 0$; the receiver’s cost of inspection is $C \in (0, B)$; both discount at $\delta \in (0, 1)$.

4 Equilibrium and Belief Updates

We focus on stationary Perfect Bayesian equilibria in the public belief state $x = (\lambda, \mu) \in [0, 1]^2$, where λ is the receiver's belief that the sender is the committed honest type and μ is the sender's belief that the receiver is the committed vigilant type. Let $\sigma(x) \in [0, 1]$ denote the strategic sender's deception probability and $\rho(x) \in [0, 1]$ the strategic receiver's inspection probability. (Committed types play **truth** and **check**, respectively.)

When $a_t^r = \text{check}$, the audit outcome (*truth* or *deceive*) is publicly disclosed; if deception is detected, termination (benchmark) or punishment (extension) follows. When $a_t^r = \text{trust}$, the sender's action remains private and no public signal is realized.

At state $x = (\lambda, \mu)$, the (public) probability of a check is

$$p^{\text{check}}(x) = \mu + (1 - \mu) \rho(x),$$

since the vigilant receiver checks with probability 1 and the strategic receiver with $\rho(x)$. Conditional on a check, the probability that the audit reveals *truth* is

$$p^{\text{T|check}}(x) = \lambda + (1 - \lambda) [1 - \sigma(x)],$$

and thus the probability of revealed deception is $(1 - \lambda)\sigma(x)$. Bayes' rule gives the updated beliefs after a check:

$$\lambda^+(x) = \frac{\lambda}{\lambda + (1 - \lambda) [1 - \sigma(x)]}, \quad \mu^+(x) = \frac{\mu}{\mu + (1 - \mu) \rho(x)}.$$

If $a_t^r = \text{trust}$, then $\mu^{\text{tr}}(x) = 0$ (the vigilant type never trusts) and $\lambda^{\text{tr}}(x) = \lambda$.

Let $V_s(x)$ and $V_r(x)$ denote the strategic sender's and receiver's stationary continuation values at $x = (\lambda, \mu)$. Given policies (σ, ρ) , the sender's Bellman equation is

$$\begin{aligned} V_s(\lambda, \mu) = \max_{\sigma \in [0, 1]} \left\{ \sigma B + \delta \left[(1 - p^{\text{check}}(x)) V_s(\lambda, 0) \right. \right. \\ \left. \left. + p^{\text{check}}(x) (1 - \sigma) V_s(\lambda^+(x), \mu^+(x)) \right] \right\}, \end{aligned} \quad (4.1)$$

since (i) deception yields B this period regardless of inspection, (ii) a check that reveals deception terminates the relationship (no continuation), and (iii) trust sends the next state to $(\lambda, 0)$. The receiver's Bellman equation is

$$\begin{aligned} V_r(\lambda, \mu) = \max_{\rho \in [0, 1]} \left\{ \rho \left(\underbrace{B p^{\text{T|check}}(x) - C}_{\text{current payoff under check}} + \delta p^{\text{T|check}}(x) V_r(\lambda^+(x), \mu^+(x)) \right) \right. \\ \left. + (1 - \rho) \left(\underbrace{B (1 - (1 - \lambda)\sigma(x))}_{\text{current payoff under trust}} + \delta V_r(\lambda, 0) \right) \right\}. \end{aligned} \quad (4.2)$$

Assumption A. (i) Per-period payoffs are bounded and do not depend on (λ, μ) except via the action-realization branches described by (4.1)–(4.2).

- (ii) The public Bayes maps satisfy $\lambda^+(\lambda)$ increasing in λ and $\mu^+(\mu, \rho)$ increasing in μ and (weakly) decreasing in ρ ; the “trust reset” branch goes to $(\lambda, 0)$.
- (iii) $p^{\text{check}}(\lambda, \mu) = \mu + (1 - \mu) \rho(\lambda, \mu)$ is (weakly) increasing in μ .
- (iv) The evaluation operator for fixed policies (σ, ρ) is a δ -contraction on bounded functions (Banach).

Proposition 4.1. *Fix any measurable stationary policies (σ, ρ) . Let $(V_s^{\sigma, \rho}, V_r^{\sigma, \rho})$ be the unique evaluation fixed point of (4.1)–(4.2). Under Assumption A,*

$$\begin{aligned} V_s^{\sigma, \rho} \text{ and } V_r^{\sigma, \rho} &\text{ are increasing in } \lambda; \\ V_s^{\sigma, \rho} &\text{ is (weakly) decreasing in } \mu; \\ V_r^{\sigma, \rho} &\text{ is (weakly) increasing in } \mu. \end{aligned}$$

Proof. Define the policy–evaluation operator $\mathcal{T}_{\sigma, \rho}$ (RHS of (4.1)–(4.2)). By (iv) it is a δ -contraction. *Monotonicity in λ .* Each branch on the RHS is monotone in λ : $p_T(\lambda, \mu) = 1 - (1 - \lambda)\sigma(\lambda, \mu)$ increases in λ ; the continuation nodes are $(\lambda, 0)$ and $(\lambda^+(\lambda), \mu^+(\mu, \rho))$ with λ^+ increasing. Hence if V_i is increasing in λ , so is $\mathcal{T}_{\sigma, \rho} V_i$; iterating from any bounded seed and taking the contraction limit preserves the order (standard monotone–operator argument). *Monotonicity in μ .* On the sender’s RHS, $p^{\text{check}}(\lambda, \mu)$ increases in μ and the survival branch following trust goes to $(\lambda, 0)$ while the check branch goes to (λ^+, μ^+) with μ^+ increasing in μ . Since higher μ raises the inspection/termination risk and (weakly) lowers the continuation via μ^+ , $\mathcal{T}_{\sigma, \rho} V_s$ is (weakly) decreasing in μ whenever V_s is. On the receiver’s RHS, higher μ raises p^{check} and μ^+ , both improving discipline; thus $\mathcal{T}_{\sigma, \rho} V_r$ is (weakly) increasing in μ whenever V_r is. Contraction again preserves these orders at the fixed point. $\square$

Assumption B. At any joint mixing state x^* , the product $G(\mu) := p^{\text{check}}(\lambda^*, \mu) \cdot [V_s(\lambda^+(\lambda^*), \mu^+(\mu, \rho)) - V_s(\lambda^*, 0)]$ is (weakly) increasing in μ .¹

Proposition 4.2. *Let $\Delta_s(\lambda, \mu) := B - \delta p^{\text{check}}(\lambda, \mu) [V_s(\lambda^+, \mu^+) - V_s(\lambda, 0)]$ and $\Delta_r(\lambda, \mu) := C - \delta [p_T(\lambda, \mu) V_r(\lambda^+, \mu^+) - V_r(\lambda, 0)]$, where $p_T(\lambda, \mu) = 1 - (1 - \lambda)\sigma(\lambda, \mu)$. Under Assumptions A–B and Theorem 5.1,*

$$\partial_\lambda \Delta_r(\lambda, \mu) < 0 \quad \text{and} \quad \partial_\mu \Delta_s(\lambda, \mu) > 0$$

whenever the partials exist; with directional derivatives otherwise. Hence the receiver’s best reply is (weakly) decreasing in λ , and the sender’s best reply is (weakly) decreasing in μ (single-crossing).

Proof. For Δ_r , $p_T(\lambda, \mu)$ is increasing in λ , $V_r(\lambda^+, \mu^+)$ is increasing in λ by Theorem 5.1 and λ^+ is increasing; thus the continuation term increases in λ , implying Δ_r is strictly decreasing

¹A sufficient, checkable condition is $\partial_\mu p^{\text{check}}(\lambda^*, \mu) \cdot J(\mu) \geq -p^{\text{check}}(\lambda^*, \mu) \cdot L_\mu \cdot \partial_\mu \mu^+(\mu, \rho)$ for all μ near μ^* , where $J(\mu) := V_s(\lambda^+(\lambda^*), \mu^+) - V_s(\lambda^*, 0) > 0$ and L_μ is a local Lipschitz bound on $-\partial_\mu V_s$. Under ρ locally independent of μ , $\partial_\mu \mu^+ = \rho/(\mu + (1 - \mu)\rho)^2$ and the inequality reduces to an elasticity bound ensuring the hazard effect dominates the continuation–dampening effect.

(unless the check branch has zero weight). For Δ_s , Assumption B says the product $p^{\text{check}} \cdot [\cdot]$ increases in μ , so $\Delta_s = B - \delta(\cdot)$ is (weakly) increasing in μ . Directional derivatives yield the same order when kinks arise at boundaries. $\square$

Definition 4.3. A stationary PBE is a tuple (V_s, V_r, σ, ρ) and belief-update rules as above such that:

- (i) V_s, V_r satisfy the Bellman equations given (σ, ρ) ;
- (ii) σ, ρ are pointwise optimal given V_s, V_r ;
- (iii) beliefs update by Bayes' rule on the equilibrium path (and via standard refinements off path).

At any public history h^t with zero probability under (σ, ρ) , beliefs are defined as limits of posteriors along vanishing trembles: take a sequence of full-support strategy profiles $\{(\sigma^\varepsilon, \rho^\varepsilon)\}_{\varepsilon \downarrow 0}$ that coincides with (σ, ρ) on-path, compute the Bayes posteriors at h^t under $(\sigma^\varepsilon, \rho^\varepsilon)$, and let (λ, μ) be any limit point as $\varepsilon \rightarrow 0$. If a *check* or *deceive* action is observed at a history where the equilibrium assigns zero probability to it, it is treated as a tremble and the same Bayes operator used on-path (for checks/truthful checks) is applied to determine (λ^+, μ^+) . Beliefs at termination (and in the finite-punishment states) are degenerate and history-independent. This convention makes the right-hand sides of (4.1)–(4.2) well defined at every public history and preserves the continuity statements recorded in Theorem 8.4.

Remark 4.4. If audit outcomes upon *check* are not publicly disclosed when they reveal *truth*, then λ does not update publicly after such checks. The public state remains Markovian with transitions: $(\lambda, \mu) \rightarrow (\lambda, 0)$ after *trust*; $(\lambda, \mu) \rightarrow (\lambda, \mu^+)$ after *check* and no detection; and absorption upon detected deception. Analysis then proceeds with public (λ, μ) but the receiver's private belief about honesty strictly dominates the public λ ; the stationary PBE structure survives under a mild public-reporting commitment (formal details omitted here).

5 Cutoff Characterization

We derive sender/receiver indifference conditions under stationary mixing and establish monotone cutoff policies. Throughout this section we work with stationary Markov strategies $\sigma(\lambda, \mu) \in [0, 1]$ and $\rho(\lambda, \mu) \in [0, 1]$, with committed types playing *truth* (sender) and *check* (receiver). Recall

$$p^{\text{check}}(x) = \mu + (1 - \mu) \rho(x), \quad p^{\text{T|check}}(x) = \lambda + (1 - \lambda) [1 - \sigma(x)],$$

and the posteriors after a check (with public disclosure) are

$$\lambda^+(x) = \frac{\lambda}{\lambda + (1 - \lambda) [1 - \sigma(x)]}, \quad \mu^+(x) = \frac{\mu}{\mu + (1 - \mu) \rho(x)}.$$

5.1 Indifference (mixing) equations

Fix a public belief state $x = (\lambda, \mu)$ where both strategic players mix. Let V_s, V_r be the stationary continuation values.

Comparing **truth** and **deceive** in the Bellman equation yields the sender's indifference:

$$B = \delta p^{\text{check}}(x) \left[V_s(\lambda^+(x), \mu^+(x)) - V_s(\lambda, 0) \right]. \quad (5.1)$$

Interpretation. The one-shot gain B from deception equals the discounted reputational dividend from a *truthful* audit: with probability p^{check} a check occurs; conditional on not being terminated, the continuation jumps from the “trust after no-check” state $(\lambda, 0)$ to the “verified-truth” state (λ^+, μ^+) .

Comparing **check** and **trust** yields the receiver's indifference:

$$C = \delta \left[p^{\text{check}}(x) V_r(\lambda^+(x), \mu^+(x)) - V_r(\lambda, 0) \right]. \quad (5.2)$$

The flow cost C must equal the discounted gain from improving beliefs when a check *does not* terminate (i.e., reveals truth). In expectation, termination branches contribute no continuation value.

Assumption C. Suppose the following holds:

- (i) Stage payoffs depend on (λ, μ) only via public survival/termination and the inspection cost; $B, C \in (0, \infty)$ and $\delta \in (0, 1)$ are fixed.
- (ii) Upon **trust**, beliefs reset to $(\lambda, 0)$; upon a truthful, publicly observed check, beliefs update to $(\lambda^+(\lambda, \mu), \mu^+(\lambda, \mu))$, where

$$\lambda^+(\lambda, \mu) = \frac{\lambda}{\lambda + (1 - \lambda)(1 - \sigma(\lambda, \mu))}, \quad \mu^+(\lambda, \mu) = \frac{\mu}{\mu + (1 - \mu)\rho(\lambda, \mu)}.$$

- (iii) $\lambda^+(\cdot, \mu)$ is increasing and concave in λ for every μ , and $\mu^+(\lambda, \cdot)$ is increasing in μ for every λ .
- (iv) The policy space $\Sigma \times P$ is the compact product of measurable maps $[0, 1]^2 \rightarrow [0, 1]$ endowed with the sup norm.

These primitives are delivered in the Online Appendix (see Proposition OA3.8: monotone best replies under value monotonicity and the public-check recursion).

Proposition 5.1. *Fix measurable policies (σ, ρ) . The policy-evaluation operator induced by (4.1)–(4.2) is a δ -contraction whose unique fixed point $(V_s^{\sigma, \rho}, V_r^{\sigma, \rho})$ satisfies:*

1. $V_s^{\sigma, \rho}$ and $V_r^{\sigma, \rho}$ are increasing in λ ;
2. $V_s^{\sigma, \rho}$ is (weakly) decreasing in μ , while $V_r^{\sigma, \rho}$ is (weakly) increasing in μ ;
3. Let $\Delta_r := V_r^{\text{check}} - V_r^{\text{trust}}$ and $\Delta_s := V_s^{\text{deceive}} - V_s^{\text{truth}}$. Then $\Delta_r(\lambda, \mu)$ is weakly decreasing in λ , and $\Delta_s(\lambda, \mu)$ is weakly decreasing in μ .

Proof. For fixed (σ, ρ) the Bellman map is affine in (V_s, V_r) with discount $\delta < 1$, hence a contraction. Order preservation: increasing λ raises (i) the probability the check finds truth and (ii) survival continuation, while increasing μ raises the check probability—hurting the sender (more termination risk) and helping the receiver (more discipline). Start value iteration from 0, apply monotone induction using the concavity of λ^+ and monotonicity of μ^+ ((Assumption C(iii))), and pass to the limit by contraction to obtain (a)–(b). The action-difference formulas in (A.1)–(A.2) together with (a)–(b) and the properties of λ^+, μ^+ imply (c). $\square$

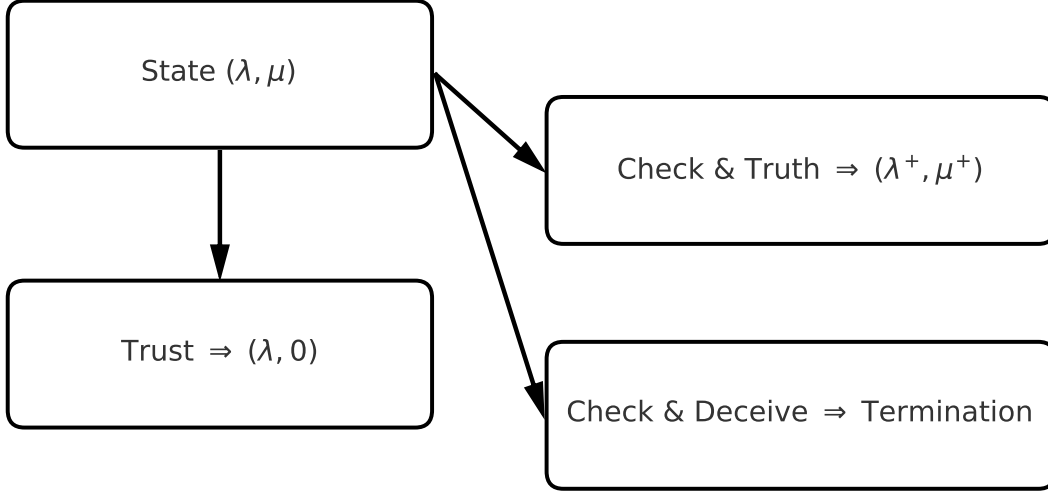


Figure 1: Public belief recursion. Trust: $(\lambda, 0)$. Check & Truth: (λ^+, μ^+) . Check & Deceive: termination.

5.2 Monotonicity and threshold structure

Lemma 5.2. *Under Assumption C, any stationary PBE admits (weakly) monotone policies: (i) $\rho(\lambda, \mu)$ is weakly decreasing in λ ; (ii) $\sigma(\lambda, \mu)$ is weakly decreasing in μ . Consequently, there exist cutoffs $\lambda^*(\mu)$ and $\mu^*(\lambda)$ with the usual threshold form.*

Proof. See [Appendix](#). □

By Theorems 4.2 and 5.1, the sign conditions required for single-crossing hold in our environment, so the receiver’s (sender’s) best reply is weakly decreasing in λ (in μ).

5.3 A simple closed-form benchmark

To provide intuition and sharp formulas, we now analyze a tractable “one-step deterrence” benchmark that captures the idea that public checks discipline near-term behavior.

Assumption D. At a mixing state $x = (\lambda, \mu)$:

- (A1) The strategic receiver mixes with a *constant* inspection rate $\rho(x) = \bar{\rho} \in (0, 1)$ locally, and the strategic sender mixes with $\sigma(x) = \bar{\sigma} \in (0, 1)$ locally.
- (A2) A check that reveals *truth* raises vigilance belief to $\mu^+ = \bar{\mu} \in (\mu, 1]$, which (via public observability of checks) reduces the next-period deception probability from $\bar{\sigma}$ to $\bar{\sigma}' \in [0, \bar{\sigma})$ for exactly one period (a deterrence window), after which policies revert to $(\bar{\sigma}, \bar{\rho})$ unless the game has absorbed.

- (A3) If the receiver *trusts*, the vigilance belief collapses to 0 next period (public trust), so the sender's continuation equals the trust-absorption value $V_s(\lambda, 0) = \frac{B}{1-\delta}$.
(A4) Termination after detected deception delivers zero continuation for both players.

Assumption D is a local reduced-form capturing that (i) checks are publicly observed and temporarily stiffen perceived vigilance; (ii) trusting publicly reveals non-vigilance.

Lemma 5.3. *Under Assumption D, the sender's indifference (5.1) pins down the total inspection probability at the mixing point as*

$$p^{\text{check}*} = \frac{1-\delta}{\delta}. \quad (5.3)$$

and, more generally, it obeys the bounds (5.7) under Assumption A'.

Equivalently, $\mu + (1-\mu)\bar{\rho} = (1-\delta)/\delta$ at the sender's cutoff. Moreover, if the receiver's one-step deterrence reduces the next-period deception probability from $\bar{\sigma}$ to $\bar{\sigma}'$, the receiver's indifference (5.2) yields a closed-form honesty cutoff

$$\lambda^* = 1 - \frac{C}{\delta B p^{\text{check}*} (\bar{\sigma} - \bar{\sigma}')} , \quad (5.4)$$

provided the RHS lies in $(0, 1)$. In particular, λ^* is decreasing in the deterrence gap $(\bar{\sigma} - \bar{\sigma}')$ and increasing in C/B .

If (V_s, V_r) are only piecewise differentiable near the mixing state, the sender and receiver indifference maps admit one-sided derivatives. If the indifference loci cross with strictly opposite one-sided slopes (a transversality condition), the joint mixing point is locally unique; non-differentiabilities arise only at policy kinks.

Proof. See Appendix. □

Proposition 5.4. *Let $J(\lambda, \mu) := V_s(\lambda^+(\lambda), \mu^+(\mu, \rho)) - V_s(\lambda, 0)$ denote the continuation jump after a truthful check at (λ, μ) . At a sender mixing state $x^* = (\lambda^*, \mu^*)$ the indifference (5.1) implies*

$$p^{\text{check}}(x^*) = \frac{B}{\delta J(x^*)}.$$

Suppose for parameters in a neighborhood of x^* we have $\underline{J} \leq J(\lambda, \mu) \leq \bar{J}$ with $0 < \underline{J} \leq \bar{J} < \infty$. Then

$$\frac{B}{\delta \bar{J}} \leq p^{\text{check}*} \leq \frac{B}{\delta \underline{J}}. \quad (5.5)$$

In particular, if $J(x^*)$ is ε -close (multiplicatively) to the perpetuity wedge $B/(1-\delta)$, i.e.

$$J(x^*) \in \left[\frac{B}{1-\delta}(1-\varepsilon), \frac{B}{1-\delta}(1+\varepsilon) \right],$$

then

$$p^{\text{check}*} \in \left[\frac{1-\delta}{\delta} \cdot \frac{1}{1+\varepsilon}, \frac{1-\delta}{\delta} \cdot \frac{1}{1-\varepsilon} \right], \quad (5.6)$$

so the closed form $p^{\text{check}*} = (1-\delta)/\delta$ is robust up to a factor $1 \pm \varepsilon$.

Proof. At x^* , sender indifference gives $p^{\text{check}*} = B/(\delta J(x^*))$. Bounding $J(x^*)$ between $\underline{J}$ and $\bar{J}$ yields (5.7). The interval (5.6) follows by substituting the ε -band around $B/(1 - \delta)$. $\square$

Remark 5.5. The perpetuity $J \simeq B/(1 - \delta)$ arises when (i) a truthful check at x^* locally resets the public state to one whose subsequent hazard path and policy pair are (approximately) stationary, and (ii) the hazard in this neighborhood is small enough that continuation risk is second order. Two convenient sufficient routes are: (a) the continuous-time bridge in Section 10.4 with discount rate β and locally constant intensities near the cutoff, which implies $J \rightarrow B/\beta$ and $\delta = e^{-\beta\Delta}$ delivers $J \simeq B/(1 - \delta)$ as $\Delta \downarrow 0$; (b) a discrete stationary neighborhood where $\lambda^+(\lambda^*) \approx \lambda^*$ and $\mu^+(\mu^*, \rho) \approx \mu^*$ (e.g., via symmetry or by construction in the punishment design), which makes the post-check continuation a geometric sum with common ratio δ . When these invariances only hold approximately, Theorem 5.6 provides the corresponding bounds on $p^{\text{check}*}$.

Assumption E. At the sender's mixing state x^* , the truthful-check continuation jump

$$W(x^*) := V_s(\lambda^+(x^*), \mu^+(x^*)) - V_s(\lambda^*, 0)$$

lies in the interval $\left[(1 - \epsilon) \frac{B}{1 - \delta}, (1 + \epsilon) \frac{B}{1 - \delta}\right]$ for some $\epsilon \in [0, 1)$.

Proposition 5.6. *Under Assumption E, the sender's indifference (5.1) implies*

$$\frac{1 - \delta}{\delta} \cdot \frac{1}{1 + \epsilon} \leq p^{\text{check}*} \leq \frac{1 - \delta}{\delta} \cdot \frac{1}{1 - \epsilon}. \quad (5.7)$$

If a truthful public check locally reduces next-period deception from $\bar{\sigma}$ to $\bar{\sigma}'$ (one-step deterrence as in Lemma 5.3), then the receiver's cutoff satisfies

$$1 - \frac{C}{\delta B p^{\text{check}*} (\bar{\sigma} - \bar{\sigma}')} \leq \lambda^* \leq 1 - \frac{C}{\delta B p^{\text{check}*} (\bar{\sigma} - \bar{\sigma}')} \cdot \frac{1 - \epsilon}{1 + \epsilon}, \quad (5.8)$$

whenever the right-hand sides lie in $(0, 1)$. Hence λ^ is increasing in C/B , decreasing in $(\bar{\sigma} - \bar{\sigma}')$, and continuous in ϵ at 0.*

Proof. From (5.1), $B = \delta p^{\text{check}*} W(x^*)$. Using $W(x^*) \in \left[(1 - \epsilon) \frac{B}{1 - \delta}, (1 + \epsilon) \frac{B}{1 - \delta}\right]$ yields (5.7). For (5.8), repeat the algebra in Lemma 5.3: the continuation gain entering the receiver's indifference scales linearly with $W(x^*)$, so the same ϵ -bounds propagate to λ^* . $\square$

At the sender's margin, the total inspection probability is pinned down purely by patience: $p^{\text{check}*} = (1 - \delta)/\delta$ decreases in δ (more patient players require rarer inspections to deter). At the receiver's margin, stronger deterrence $(\bar{\sigma} - \bar{\sigma}')$ or cheaper inspection C lowers the honesty cutoff λ^* , expanding the trust region; larger B (higher temptation) raises the benefit of checking and thus also lowers λ^* via (5.4).

Remark 5.7. Lemma 5.3 provides transparent, closed-form intuition; it abstracts from multi-period belief feedback beyond one step. In the full model without Assumption D, (5.1)–(5.2) jointly determine functions $\lambda^*(\mu)$ and $\mu^*(\lambda)$. The qualitative monotonicity (Lemma 5.2) and the signs of comparative statics survive in general. In the closed-form benchmark, λ^* increases in B , decreases in C , and decreases in δ .

5.4 Punishment extension

We evaluate planner welfare as the receiver’s surplus from truthful trade net of inspection costs, treating the sender’s temptation benefit and any penalties as transfers.² Let $W(\lambda_0, \mu_0)$ denote expected discounted welfare from prior (λ_0, μ_0) .

Theorem 5.8. *Along any stationary policy pair (σ, ρ) starting at (λ_0, μ_0) ,*

$$W(\lambda_0, \mu_0) = \mathbb{E} \left[\sum_{t \geq 0} \delta^t \left\{ B \cdot p_T(x_t) - C \cdot p^{\text{check}}(x_t) \right\} \right], \quad p_T(\lambda, \mu) = 1 - (1 - \lambda)\sigma(\lambda, \mu), \quad (5.9)$$

where $x_t = (\lambda_t, \mu_t)$ is the public state and $p^{\text{check}}(\lambda, \mu) = \mu + (1 - \mu)\rho(\lambda, \mu)$. Equivalently, writing the termination hazard $h(\lambda) = p^{\text{check}}(\lambda)(1 - \lambda)\sigma(\lambda)$ as in (8.2),

$$W(\lambda_0, \mu_0) = \mathbb{E} \left[\sum_{t \geq 0} \delta^t \left\{ B \cdot (1 - h(\lambda_t) - p^{\text{check}}(\lambda_t)\lambda_t\sigma(\lambda_t)) - C \cdot p^{\text{check}}(\lambda_t) \right\} \right]. \quad (5.10)$$

Proof. Welfare each period is (surplus from truthful trade) minus (inspection cost). Truth at x_t occurs unless deception and detection coincide; hence its probability is $p_T(x_t)$, yielding (5.9). The hazard identity $h(\lambda) = p^{\text{check}}(\lambda)(1 - \lambda)\sigma(\lambda)$ and $p_T = 1 - (1 - \lambda)\sigma$ produce (5.10) by algebra. Linearity of expectation gives the result. $\square$

Corollary 5.9. (i) *Holding policies fixed, $\partial W / \partial C = -\mathbb{E}[\sum_{t \geq 0} \delta^t p^{\text{check}}(x_t)] < 0$ and $\partial W / \partial B = \mathbb{E}[\sum_{t \geq 0} \delta^t p_T(x_t)] > 0$. (ii) *In equilibrium, higher transparency (silent-audit disclosure q or precision κ) weakly raises W by reducing the deception region and (for fixed λ) weakly lowering the hazard $h(\lambda)$; see Theorem 10.1 and §10.2. (iii) *When the same parameter B governs both temptation and truthful surplus (our baseline normalization), the net effect on W is a priori ambiguous because σ, ρ adjust endogenously; the sign is positive when the induced change in hazard is second order (small) relative to the direct surplus term in (5.9).***

Remark 5.10. Equation (5.10) shows that front-loading inspections is attractive when B/C is large: early checks lower $h(\lambda)$ by pushing λ up (via λ^+) and reduce future p^{check} along the tapering path. Conversely, inspection spurts temporarily raise $h(\lambda)$ mechanically, generating bunching of terminations—our empirical prediction in §9.

6 Punishment extension

The fixed-point existence extends to finite Markov punishment with an augmented compact state; the contraction and upper-hemicontinuity steps carry over (see the Online Appendix).

We now replace termination upon detected deception with a *publicly observed punishment phase*. Upon a period- t history with (deceive, check) revealed by an audit, the game transitions

²This matches our applications (auditing/certification) where value comes from compliant trade and inspections are resource costs; sender temptation rents are not social surplus.

to a punishment state for a finite number of periods $T \in \mathbb{N}$ and then returns to the normal phase. Punishment policies are public and history-dependent; they may prescribe, for instance, automatic inspections and exclusion from trust.

Assumption F. The augmented public state is $\tilde{x} = (\lambda, \mu, s)$ with $(\lambda, \mu) \in [0, 1]^2$ and $s \in \mathcal{S}$, where $\mathcal{S} = \{\mathbf{N}\} \cup \{\mathbf{P}_1, \dots, \mathbf{P}_L\} \cup \{\mathbf{T}\}$ consists of the normal mode $\mathbf{N}$, a finite punishment chain of length $L \in \mathbb{N}$, and an absorbing termination state $\mathbf{T}$.

- (i) Payoffs are bounded and depend on $\tilde{x}$ only through (λ, μ) and s ; discount $\delta \in (0, 1)$.
- (ii) In mode $\mathbf{N}$, the one-period transitions are as in the baseline (public check or trust), except that (**deceive**, **check**) moves to $\mathbf{P}_1$ (or $\mathbf{T}$ in the benchmark).
- (iii) In punishment, $\mathbf{P}_k$ moves to $\mathbf{P}_{k+1}$ for $k < L$ and to $\mathbf{N}$ from $\mathbf{P}_L$ (or to $\mathbf{T}$ if specified), with transition probabilities and posteriors given by continuous functions of (λ, μ) and the policy pair (σ, ρ) .
- (iv) All belief updates $\tilde{x} \mapsto (\lambda^+, \mu^+, s')$ are continuous in (λ, μ) and measurable in s , with $\mathbf{T}$ absorbing.

Lemma 6.1. *Under Assumption F, the policy-evaluation operator associated with (6.1)–(6.2) is a δ -contraction on the bounded functions over $[0, 1]^2 \times \mathcal{S}$ and maps continuous functions into continuous functions. Therefore, a stationary PBE exists; its value functions (V_s, V_r) are the unique bounded fixed points and are continuous in (λ, μ) for each s , with piecewise continuity across modes and continuity at $\mathbf{N} \leftrightarrow \mathbf{P}_1$ and $\mathbf{P}_L \leftrightarrow \mathbf{N}$ implied by (iii)–(iv).*

Proof. $\mathcal{S}$ is finite, so $[0, 1]^2 \times \mathcal{S}$ is compact. Fix measurable (σ, ρ) . The Bellman right-hand sides in (6.1)–(6.2) are affine in (V_s, V_r) with coefficient $\delta < 1$, hence define a δ -contraction on the sup-norm space of bounded functions over the augmented state. Continuity of the policy-evaluation operator follows because all branch probabilities and update maps are continuous in (λ, μ) by (iii)–(iv) and s takes finitely many values. Uniqueness of $(V_s^{\sigma, \rho}, V_r^{\sigma, \rho})$ follows by Banach; upper hemicontinuity/convexity of best replies is as in the baseline (affine in own mixed action, continuous in state/values). Kakutani–Fan–Glicksberg then gives a stationary fixed point (σ^*, ρ^*) . $\square$

Let $V_s^{\text{pun}}(x)$ and $V_r^{\text{pun}}(x)$ denote the strategic sender’s and receiver’s continuation values at the *entry* into punishment after detection at public state $x = (\lambda, \mu)$. To keep the extension parsimonious, we impose a *simple* punishment: for T periods the receiver *must* inspect (so trade occurs only under verified truth), and at the end of the T periods the public state resets to $(\lambda, 0)$ (the same trust-absorption state as under **trust**). Under this scheme,

$$V_s^{\text{pun}}(x) = \delta^T V_s(\lambda, 0),$$

$$V_r^{\text{pun}}(x) = \sum_{k=0}^{T-1} \delta^k \underbrace{[B \cdot p_{\text{pun}}^{\text{T|check}} - C]}_{\text{per-period payoff in punishment}} + \delta^T V_r(\lambda, 0),$$

where $p_{\text{pun}}^{\text{T|check}} \in [0, 1]$ is the probability a check in punishment reveals **truth** (equal to 1 against the committed honest sender and to $1 - \sigma$ against the strategic sender if he mixes). In the *maximal* punishment we take $p_{\text{pun}}^{\text{T|check}} = 1$ only when the sender is the honest type (the

strategic sender's deception under mandatory checks brings no surplus and is immediately re-detected).

Let V_s, V_r be the values in the normal phase. Relative to Section 4, only the branches following detected deception change (they now lead to $V_s^{\text{pun}}, V_r^{\text{pun}}$ instead of 0):

$$\begin{aligned} V_s(\lambda, \mu) = \max_{\sigma \in [0,1]} \left\{ \sigma \left[B + \delta \left(p^{\text{check}}(x) V_s^{\text{pun}}(x) + (1 - p^{\text{check}}(x)) V_s(\lambda, 0) \right) \right] \right. \\ \left. + (1 - \sigma) \delta \left(p^{\text{check}}(x) V_s(\lambda^+(x), \mu^+(x)) \right. \right. \\ \left. \left. + (1 - p^{\text{check}}(x)) V_s(\lambda, 0) \right) \right\}, \end{aligned} \quad (6.1)$$

$$\begin{aligned} V_r(\lambda, \mu) = \max_{\rho \in [0,1]} \left\{ \rho \left(\overbrace{B p^{\text{T|check}}(x) - C}^{\text{current payoff under check}} + \delta \left[p^{\text{T|check}}(x) V_r(\lambda^+(x), \mu^+(x)) \right. \right. \right. \\ \left. \left. + (1 - p^{\text{T|check}}(x)) V_r^{\text{pun}}(x) \right] \right) \\ \left. + (1 - \rho) \left(\underbrace{B [1 - (1 - \lambda)\sigma(x)]}_{\text{current payoff under trust}} + \delta V_r(\lambda, 0) \right) \right\}. \end{aligned} \quad (6.2)$$

The punishment mode set $\mathcal{S} = \{\mathbf{N}, \mathbf{P}_1, \dots, \mathbf{P}_L\}$ is finite; transition probabilities across modes are continuous in (λ, μ) and independent of past beyond the current mode; and within each mode the Bayes maps (λ^+, μ^+) satisfy Assumption A(ii) with the same bounded per-period payoffs.

Lemma 6.2. *Under Assumption F, the augmented state space $[0, 1]^2 \times \mathcal{S}$ is compact, and for any stationary policies (σ, ρ) the evaluation operator on bounded functions is a δ -contraction mode-by-mode. Hence (V_s, V_r) exist and are continuous in (λ, μ) in every mode. Moreover, the stationary PBE existence result in Theorem 8.3 extends verbatim to the punishment model via Kakutani–Fan–Glicksberg.*

Proof. Fix (σ, ρ) . With $\mathcal{S}$ finite and transitions continuous, the right-hand sides of (6.1)–(6.2) define a δ -contraction on $\ell_\infty([0, 1]^2 \times \mathcal{S})$, yielding a unique fixed point that is continuous by standard parametric contraction arguments (Berge). Best-reply correspondences are convex-valued and upper hemicontinuous as in Lemma 7.1, so Kakutani gives a stationary fixed point. $\square$

At a sender cutoff $\mu^*(\lambda)$ and a receiver cutoff $\lambda^*(\mu)$ where both strategic players mix,

$$\text{S:} \quad B = \delta p^{\text{check}}(x) \left[V_s(\lambda^+(x), \mu^+(x)) - V_s^{\text{pun}}(x) \right], \quad (6.3)$$

$$\text{R:} \quad C = \delta \left[p^{\text{T|check}}(x) V_r(\lambda^+(x), \mu^+(x)) + (1 - p^{\text{T|check}}(x)) V_r^{\text{pun}}(x) - V_r(\lambda, 0) \right]. \quad (6.4)$$

Compare (6.3)–(6.4) to their termination counterparts

$$\begin{aligned} B &= \delta p^{\text{check}}(x) \left[V_s(\lambda^+(x), \mu^+(x)) - V_s(\lambda, 0) \right], \\ C &= \delta \left[p^{\text{T|check}}(x) V_r(\lambda^+(x), \mu^+(x)) - V_r(\lambda, 0) \right]. \end{aligned}$$

6.1 Outcome-equivalence with finite punishment

We now show that, by choosing the punishment length T (and, if desired, its within-phase inspection intensity), one can *implement the same cutoffs* as in the termination benchmark.

Proposition 6.3. *Fix a stationary PBE of the termination benchmark with mixing at $x^* = (\lambda^*, \mu^*)$. Under Assumption F, for any termination equilibrium there exists a finite-length punishment specification (choice of L and branch probabilities in (iii)) that implements the same inspection and deception cutoffs; conversely, any such punishment equilibrium is outcome-equivalent to a termination equilibrium.*

Proof. Under the simple public punishment (mandatory checks for T periods, then reset to $(\lambda^*, 0)$), the punishment values are

$$\begin{aligned} V_s^{\text{pun}}(T) &= \delta^T V_s(\lambda^*, 0), \\ V_r^{\text{pun}}(T) &= \sum_{k=0}^{T-1} \delta^k (-C) + \delta^T V_r(\lambda^*, 0) = -C \frac{1 - \delta^T}{1 - \delta} + \delta^T V_r(\lambda^*, 0). \end{aligned}$$

In the sender indifference (6.3), the RHS is strictly increasing in T (since $V_s^{\text{pun}}(T)$ decreases in T), matching the termination RHS at $T = 0$ and approaching the “zero continuation” RHS as $T \rightarrow \infty$. By continuity there exists T_s that restores the sender equality at (λ^*, μ^*) for the fixed mixing intensities at x^* .

In the receiver indifference (6.4), the RHS is strictly decreasing in T , likewise matching termination at $T = 0$ and approaching the limit with $V_r^{\text{pun}} = -C/(1 - \delta)$ as $T \rightarrow \infty$. Hence there exists T_r that restores the receiver equality.

Because the indifference residuals are continuous in the two strategic mixing probabilities at x^* (which determine p^{check} and $p^{\text{T|check}}$), we can jointly adjust the two mixes and pick a *common* finite T so that both indifferences hold at the same (λ^*, μ^*) .

Sequential rationality of punishment (mandatory checks) follows for δ large: deviating saves C now but lowers the public reset value, which strictly reduces continuation for the receiver; similar logic applies to the sender. Therefore the punishment equilibrium implements the termination cutoffs. $\square$

Proposition 6.3 shows that one need not rely on absorption to obtain the cutoff structure: a finite, publicly observed punishment can replicate the *normal-phase* indifference conditions that pin down inspection and deception thresholds. In applications, this corresponds to “probation” or “enhanced supervision” for a finite horizon after a detected violation, after which the relationship resumes. The equivalence is constructive and uses only the receiver’s publicly visible inspections; no transfers are required.

Remark 6.4. If one allows additional instruments during punishment (e.g., exclusion from trade, capped frequency of interaction), V_s^{pun} and V_r^{pun} can be shaped more flexibly, easing the existence of a common T that implements the termination cutoffs exactly.

7 Existence and (Local) Uniqueness

We provide conditions under which a stationary Perfect Bayesian equilibrium (PBE) exists in the belief state $x = (\lambda, \mu) \in [0, 1]^2$ and show a local uniqueness result for the mixing cutoffs that solve the indifference equations (5.1)–(5.2) (termination) and their punishment analogues (6.3)–(6.4).

7.1 Assumptions

- (E1) **Primitives.** $0 < C < B < \infty$ and $\delta \in (0, 1)$. Committed types exist with independent priors $\lambda_0, \mu_0 \in (0, 1)$. Payoffs are bounded and actions are finite.
- (E2) **Public disclosure (normal phase).** We adopt the public-disclosure convention of Section 4: under $a^r = \text{check}$, the audit outcome (truth/deceive) is publicly observed; under $a^r = \text{trust}$, no public signal is realized.
- (E3) **Stationary Markov strategies and belief recursion.** Strategic policies are Borel-measurable maps $\sigma, \rho : [0, 1]^2 \rightarrow [0, 1]$. Beliefs update by Bayes' rule along the equilibrium path as in Section 4; under **trust**, $\mu' = 0$ and $\lambda' = \lambda$; under **check**, $(\lambda', \mu') = (\lambda^+(x), \mu^+(x))$ unless detection triggers absorption (termination) or the punishment state (extension). The induced transition on $[0, 1]^2$ (or on the product with a finite punishment flag) is well-defined and continuous except at the absorbing/punishment boundary.
- (E4) **Value regularity & single-crossing.** For any stationary (σ, ρ) , the Bellman equations in Section 4 (and their punishment counterparts in Section 6) admit bounded solutions (V_s, V_r) that are continuous in (λ, μ) , (weakly) increasing in λ , and such that the sender's best reply is (weakly) decreasing in μ while the receiver's is (weakly) decreasing in λ (cf. Lemma 5.2).³

7.2 Fixed-point formulation

Let $\mathcal{X} = [0, 1]^2$ and let $\mathcal{K} = \{(\sigma, \rho) : \mathcal{X} \rightarrow [0, 1]^2 \text{ Borel, bounded}\}$ endowed with the sup-norm. Given $(\sigma, \rho) \in \mathcal{K}$, define the value operators by the Bellman equations in Section 4 (termination) or Section 6 (punishment), yielding (V_s, V_r) . Define best-reply correspondences

$$\mathcal{B}_s(\sigma, \rho) = \arg \max_{\tilde{\sigma} \in [0, 1]} (\text{sender's Bellman at each } x),$$

$$\mathcal{B}_r(\sigma, \rho) = \arg \max_{\tilde{\rho} \in [0, 1]} (\text{receiver's Bellman at each } x).$$

³These monotonicities are standard for reputation models with public monitoring on the receiver side; they can be verified by induction on discounted continuation values and the fact that higher λ raises revealed-truth probabilities while higher μ raises effective inspection.

A stationary PBE corresponds to a fixed point $(\sigma^*, \rho^*) \in \mathcal{K}$ such that $\sigma^*(x) \in \mathcal{B}_s(\sigma^*, \rho^*)(x)$ and $\rho^*(x) \in \mathcal{B}_r(\sigma^*, \rho^*)(x)$ for all x , with beliefs updated as in (E3).

Lemma 7.1. *Let Φ map a policy pair (σ, ρ) to the set of stationary best replies computed at the unique evaluation fixed point of $\mathcal{T}_{\sigma, \rho}$. Then Φ has nonempty, convex values, is upper hemicontinuous on the compact, convex policy space (sup norm), and has a closed graph. Consequently, Kakutani–Fan–Glicksberg applies, yielding a stationary fixed point.⁴*

Proof. Nonemptiness/convexity follow because one-shot differences are affine in own mixed action; upper hemicontinuity and closed graph follow from continuity of the policy-to-value map (parametric contraction) and Berge’s maximum theorem. Kakutani–Fan–Glicksberg gives a fixed point. $\square$

7.3 Local uniqueness of mixing cutoffs

Let $x^* = (\lambda^*, \mu^*)$ be a mixing state where the indifference equations (5.1)–(5.2) hold under stationary (σ^*, ρ^*) .

Assumption G. At the joint mixing state $x^* = (\lambda^*, \mu^*) \in (0, 1)^2$, the action–difference maps

$$\begin{aligned} F_s(\lambda, \mu) &:= B - \delta p^{\text{check}}(\lambda, \mu) [V_s(\lambda^+, \mu^+) - V_s(\lambda, 0)], \\ F_r(\lambda, \mu) &:= C - \delta [p_T(\lambda, \mu) V_r(\lambda^+, \mu^+) - V_r(\lambda, 0)], \end{aligned}$$

with $p_T(\lambda, \mu) := \lambda + (1 - \lambda)(1 - \sigma(\lambda, \mu))$, are locally Lipschitz and directionally differentiable. Their Clarke generalized Jacobian $\partial F(x^*)$ is nonempty and contains at least one matrix J with $\det J \neq 0$.

Lemma 7.2. *Under Assumption U, there exists a neighborhood $\mathcal{N}$ of x^* on which the system $F_s(\lambda, \mu) = 0$ and $F_r(\lambda, \mu) = 0$ admits a unique solution. Hence the stationary PBE mixing cutoffs are locally unique.*

Proof. By Rademacher’s theorem, $F = (F_s, F_r)$ is a.e. differentiable and locally Lipschitz; by Clarke’s inverse function theorem, if some $J \in \partial F(x^*)$ is nonsingular, then F is locally metrically regular and admits a single-valued Lipschitz inverse on a neighborhood of $F(x^*) = (0, 0)$. Therefore the zero set of F is a singleton in a neighborhood of x^* . $\square$

Corollary 7.3. *Suppose F_s, F_r are C^1 near x^* and*

$$\partial_\mu F_s(x^*) < 0, \quad \partial_\lambda F_r(x^*) < 0,$$

while the cross partials are finite. Then the classical Jacobian

$$J = \begin{pmatrix} \partial_\lambda F_s & \partial_\mu F_s \\ \partial_\lambda F_r & \partial_\mu F_r \end{pmatrix}$$

is nonsingular and the mixing solution is locally unique. A sufficient condition for $\partial_\mu F_s < 0$ and $\partial_\lambda F_r < 0$ is the (strict) single-crossing established by Theorem 5.1 and the monotonicity of p^{check} in μ and of p_T in λ .

⁴Existence and continuity of values are stated formally in Theorem 8.3.

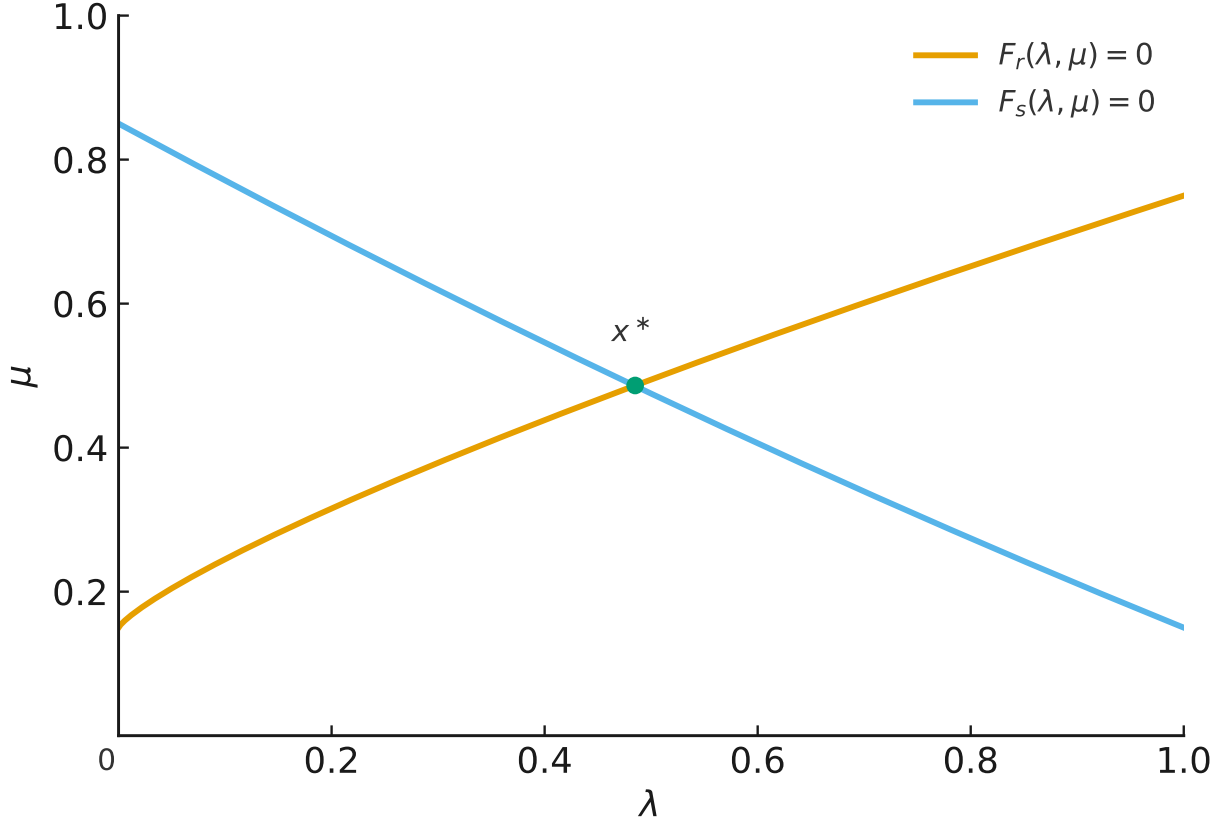


Figure 2: Receiver ($F_r(\lambda, \mu) = 0$) and sender ($F_s(\lambda, \mu) = 0$) indifference loci with unique crossing x^* .

Remark 7.4. If (V_s, V_r) are only piecewise C^1 , the indifference maps admit one-sided derivatives at x^* . If the sender and receiver indifference loci cross with strictly opposite one-sided slopes at x^* (transversality), then the joint mixing point is locally unique by the intermediate value argument applied on each side of the kink (a Clarke-type selection is not needed here).

Remark 7.5. Local uniqueness at x^* does not preclude multiple stationary PBE cutoffs globally. Multiplicity can arise if the sender and receiver indifference loci $F_s(\lambda, \mu) = 0$ and $F_r(\lambda, \mu) = 0$ intersect more than once (e.g., due to non-monotone hazards or highly curved Bayes maps far from x^*). Two forces push against multiplicity: (i) *single-crossing* (Theorem 5.1) makes F_s (weakly) downward sloping in μ and F_r (weakly) upward sloping in λ ; (ii) *monotone hazards* $h(\lambda) = p^{\text{check}}(\lambda)(1 - \lambda)\sigma(\lambda)$ that fall in λ limit additional crossings as reputations accumulate. Failures typically occur on knife-edges where one locus becomes locally flat or tangential to the other ($\det J = 0$), in which case small parameter perturbations restore transversality and single-crossing of the loci near the economically relevant region. In applications, a natural selection is the limit of stationary equilibria under vanishing payoff trembles or transparency noise (silent-audit $q \downarrow 0$), which selects the transversal crossing nearest the prior.

The locus $F_r(\lambda, \mu) = 0$ is the receiver's indifference set (check vs. trust) and $F_s(\lambda, \mu) = 0$ is the sender's indifference set (deceive vs. truth). Their transversal crossing at x^* is the

joint mixing state; by Theorems 7.2 and 7.3 this transversality delivers local uniqueness. The regions around x^* inherit the monotone best-reply structure: above the sender curve $F_s < 0$ (truth preferred), below it $F_s > 0$ (deception preferred); to the left of the receiver curve $F_r > 0$ (check preferred), to the right $F_r < 0$ (trust preferred). The plot is schematic (not to scale); comparative statics move these loci in the directions identified by Theorems 5.6, 9.1 and 10.1.

Knife-edge multiplicity corresponds to tangencies of $F_s = 0$ and $F_r = 0$ (i.e., $\det J = 0$); small parameter perturbations restore transversality and local uniqueness.

The derivatives entering DF inherit continuity from the Bellman operators and Bayes updates; nondegeneracy typically holds except at knife-edge parameter values (e.g., δ near $1/2$ in the closed-form benchmark of Lemma 5.3). The punishment extension replaces $V_s(\lambda, 0)$ in F_1 by $V_s^{\text{pun}}(x)$ and adds $V_r^{\text{pun}}(x)$ inside F_2 , but the same argument applies provided the punishment value maps are smooth in x (cf. Proposition 6.3).

8 Comparative Statics and Hazard

We study how inspection/deception cutoffs and the hazard of termination respond to (B, C, δ) .

8.1 General monotone comparative statics

Let $x = (\lambda, \mu)$ and recall the indifference equations (5.1)–(5.2). Under the single-crossing conditions of Lemma 5.2 and the regularity in Assumption (E4), totally differentiating the system $F_1(x; B, C, \delta) = F_2(x; B, C, \delta) = 0$ (see Lemma 7.2) implies:

Proposition 8.1. *Along any locally unique solution (λ^*, μ^*) to (5.1)–(5.2):*

1. **Receiver's honesty cutoff.** $\frac{\partial \lambda^*}{\partial B} > 0$, $\frac{\partial \lambda^*}{\partial C} < 0$, and (for sufficiently smooth values) $\frac{\partial \lambda^*}{\partial \delta} \leq 0$ with sign matching the net effect of δ on the receiver's discounted value of information.⁵
2. **Sender's vigilance cutoff.** $\frac{\partial \mu^*}{\partial B} > 0$, $\frac{\partial \mu^*}{\partial C} > 0$, and $\frac{\partial \mu^*}{\partial \delta} < 0$.

Intuition. A larger B makes deception more tempting; the receiver responds by checking over a *larger* honesty region (higher λ^*), and the sender needs *more* perceived vigilance to be deterred (higher μ^*). A higher C makes checking costly, so the receiver trims the checking region (lower λ^*), while the sender requires *less* vigilance to risk deception (higher μ^*). Greater patience δ lowers the sender's cutoff (less inspection suffices to deter), while the effect on the receiver depends on how δ scales the value of information relative to current costs.

⁵In the one-step benchmark below, $\partial \lambda^* / \partial \delta < 0$. In the full model, δ affects both the informational gain from checks and the continuation values, so the sign can be ambiguous without additional structure.

8.2 Closed-form benchmark

Under Assumption D (Section 5.3), Lemma 5.3 yielded

$$p^{\text{check}*} = \frac{1 - \delta}{\delta}, \quad \lambda^* = 1 - \frac{C}{(1 - \delta) B (\bar{\sigma} - \bar{\sigma}')}.$$

Hence:

$$\frac{\partial \lambda^*}{\partial B} > 0, \quad \frac{\partial \lambda^*}{\partial C} < 0, \quad \frac{\partial \lambda^*}{\partial \delta} < 0, \quad \frac{\partial p^{\text{check}*}}{\partial \delta} < 0. \quad (8.1)$$

In words: larger B or smaller C expands the checking region (raises λ^*), and greater patience reduces the total inspection frequency needed to deter (lowers $p^{\text{check}*}$) and, in this benchmark, also lowers λ^* .

8.3 Hazard of termination and welfare

Define the *per-period hazard of termination* at state x by

$$h(x) := p^{\text{check}}(x) \cdot (1 - \lambda) \cdot \sigma(x), \quad (8.2)$$

i.e., the probability that (i) a check occurs, (ii) the sender is strategic (prob. $1 - \lambda$), and (iii) the strategic sender deceives. At a mixing state $x^* = (\lambda^*, \mu^*)$, the hazard is $h^* = p^{\text{check}}(x^*) (1 - \lambda^*) \sigma(x^*)$.

Corollary 8.2. *Under the monotonicity in Lemma 5.2 and Proposition 8.1, the hazard responds as follows in a neighborhood of x^* :*

1. *Increasing B (holding fixed mixing intensity at x^*) raises h^* via a higher deception tendency and a larger checking region (higher λ^*), with the net effect typically positive.*
2. *Increasing C raises h^* by shrinking the checking region (lower λ^*) and thus increasing $(1 - \lambda^*)$; it can also weaken deterrence (higher μ^*), increasing $\sigma(x^*)$.*
3. *Increasing δ lowers h^* through a lower p^{check} requirement at the sender margin (deterrence with fewer checks); in the closed-form benchmark, this dominates and h^* falls in δ .*

8.4 Existence and continuity of stationary PBE

Theorem 8.3. *In the termination benchmark with public checks and $B, C \in (0, \infty)$, $\delta \in (0, 1)$, there exists a stationary PBE (σ^*, ρ^*) on the compact state space $[0, 1]^2$. The associated value functions (V_s, V_r) are the unique bounded fixed point of the policy-evaluation operator at (σ^*, ρ^*) and are continuous on $[0, 1]^2$, including the boundary $\mu = 0$ and absorbing/termination nodes.*

A Borel selector exists by the (Kuratowski and Ryll-Nardzewski, 1965) theorem.

Proof. Step 1. Fix measurable (σ, ρ) . The Bellman right-hand sides in (4.1)–(4.2) are affine in (V_s, V_r) with discount $\delta < 1$, hence define a δ -contraction on the Banach space of bounded measurable pairs (sup norm). Banach’s fixed point theorem yields a unique $(V_s^{\sigma, \rho}, V_r^{\sigma, \rho})$.

Step 2. All branch arguments are continuous in $x = (\lambda, \mu)$: the “trust reset” uses $(\lambda, 0)$; the truthful-check update (λ^+, μ^+) is continuous on $[0, 1]^2$ by the explicit Bayes maps; termination is an absorbing constant. The policy-evaluation operator maps $C([0, 1]^2)$ into itself and is a contraction there, so its unique fixed point is continuous (parametric contraction).

Step 3. For each x , the sender and receiver one-shot differences are continuous in x and affine in own mixed action. The best-reply correspondences (mix at indifference; choose the unique strict best action otherwise) are nonempty, convex-valued, and upper hemicontinuous (Berge’s maximum theorem). On the compact convex set of measurable policies (sup norm), the policy-to-value map $(\sigma, \rho) \mapsto (V_s^{\sigma, \rho}, V_r^{\sigma, \rho})$ is continuous (parametric contraction), hence the composed best-reply correspondence has a closed graph. Kakutani–Fan–Glicksberg yields a fixed point (σ^*, ρ^*) , which is a stationary PBE by one-step deviation. Continuity at $\mu = 0$ and on absorbing branches follows from Step 2. $\square$

Remark 8.4. The value functions (V_s, V_r) are continuous at the boundaries and absorbing nodes used in (E2)–(E5):

- (i) at $\mu = 0$, the “trust reset” branch evaluates $V_s(\lambda, 0)$ and $V_r(\lambda, 0)$, which are continuous in λ by Theorem 8.3;
- (ii) at $\lambda \in \{0, 1\}$, the Bayes maps $\lambda \mapsto \lambda^+$ are well defined as limits (concavity and monotonicity in Theorem 5.1 yield continuity at the endpoints);
- (iii) termination is an absorbing constant, hence continuous;
- (iv) in the finite-punishment extension, the augmented state $[0, 1]^2 \times \mathcal{S}$ is compact and transitions across $\mathbf{N} \leftrightarrow \mathbf{P}_1$ and $\mathbf{P}_L \leftrightarrow \mathbf{N}$ are continuous by Assumption F, so the evaluation fixed point remains continuous mode-by-mode and at mode switches.

Off-path updates at zero-probability histories are defined by limits of nearby mixed strategies, which preserves continuity of the right-hand sides of (4.1)–(4.2).

Whenever Bayes’ rule applies, beliefs are updated by Bayes. If a check occurs at a public state x where $\rho(x) = 0$ in equilibrium (or an action is otherwise off path), beliefs are the limits of posteriors under nearby fully mixed profiles. Formally, take sequences $\{\rho_\varepsilon\}$ and $\{\sigma_\varepsilon\}$ with $\rho_\varepsilon(x) > 0$, $\sigma_\varepsilon(x) \in (0, 1)$ and $\rho_\varepsilon \rightarrow \rho$, $\sigma_\varepsilon \rightarrow \sigma$ uniformly; then set the posteriors (λ^+, μ^+) at x to $\lim_{\varepsilon \downarrow 0} (\lambda_\varepsilon^+, \mu_\varepsilon^+)$. This ensures sequential rationality and well-defined updates at zero-probability histories used in (4.1)–(4.2).

The per-period *stage* surplus equals $B - C \cdot p^{\text{check}}(x)$ (checks waste C regardless of the sender’s type), while termination destroys continuation value with hazard $h(x)$. Thus, for any policy pair (σ, ρ) and belief path,

$$\begin{aligned} \text{Welfare}(x) &= \mathbb{E} \left[\sum_{t \geq 0} \delta^t (B - C \cdot p^{\text{check}}(x_t)) \right] \\ &\quad - (\text{expected continuation loss from termination events}). \end{aligned}$$

Hence welfare is decreasing in C (direct cost and higher hazard) and increasing in δ (more weight on future surplus and stronger deterrence with fewer checks). The effect of B is ambiguous: higher B raises the pie but also raises deception incentives and the hazard; in the closed-form benchmark, the net effect is positive if $(1 - \delta)$ is small and $(\bar{\sigma} - \bar{\sigma}')$ is large.

We obtain three testable predictions. First, inspection intensity should decline as the sender’s honesty reputation improves. In data, this shows up as fewer audits, reviews, or verifications following longer streaks without detected deception, holding observables fixed. A simple way to document this is to regress the probability of inspection on lagged public reputation (or its sufficient statistic, such as the share of periods with clean outcomes), or to estimate a policy function where the inspection rate is monotone in accumulated honesty.

Second, terminations should bunch after inspection spurts. Whenever inspection activity rises—because of a policy change, an enforcement campaign, or a shock to monitoring capacity—the termination hazard mechanically increases and then decays as the stock of honesty rebuilds. Event-study plots around exogenous increases in inspections should therefore display a short-run spike in exits or sanctions followed by a decline, with stronger effects in markets where deception is more tempting.

Third, environments with a higher temptation-to-cost ratio (higher benefits from deception relative to checking costs) should feature more front-loaded inspections and higher early termination hazards. Cross-sectionally, this can be examined by comparing sectors or platforms with higher expected gains from cheating or lower verification costs; within a setting, changes in fees, penalties, or detection technology that shift this ratio should move inspection timing toward the front of relationships and raise near-term exit rates. Together these implications map directly to auditing, certification, and platform-governance data, where inspection choices, public signals, and exits are routinely recorded.

9 Applications and Testable Predictions

The framework maps naturally to environments in which inspections are publicly observed while violations are hidden unless inspected.

Supervisors announce or visibly conduct audits; misreporting is unobserved unless an audit occurs. Our model interprets inspection intensity as a reputation device that disciplines hidden manipulation. It predicts front-loaded audits when B/C is high and a tapering schedule as honesty reputation accumulates.

Standards bodies (SROs, certification platforms) conduct visible conformity checks; non-compliance is detected only when checked. Visible enforcement builds a reputation for vigilance, sustaining high compliance with fewer checks over time.

Patrols and inspections are publicly observable; illegal acts or quality lapses remain hidden unless inspected. Public inspection intensity shapes offenders’ beliefs and deters hidden violations; publicity about recent inspections creates short deterrence windows.

Insurers run spot audits on claims; audit policies are often public or inferable. Higher expected gains from fraud (large claims) justify higher early audit intensity, declining with a claimant/provider’s accumulated honesty record.

Platforms publicize moderation sweeps or verification campaigns; bad behavior is observed only if caught. Visible campaigns reduce deceptive behavior temporarily; sustaining deterrence requires a baseline inspection rate.

Let λ_t be a (proxy) for accumulated honesty reputation and p_t^{check} the observed inspection intensity.

1. *Monotone inspection.* p_t^{check} is (weakly) decreasing in λ_t (Lemma 5.2); inspection intensity tapers with accumulated honest history.
2. *Hazard bunching.* Terminations spike following inspection spurts: the period hazard $h_t = p_t^{\text{check}} \cdot (1 - \lambda_t) \cdot \sigma_t$ rises mechanically with p_t^{check} and then falls as λ_t increases (Corollary 8.2).
3. *B/C shifts.* Environments or episodes with higher B/C (larger temptation, cheaper detection) exhibit higher early p_t^{check} and higher λ^* (Proposition 8.1); technological cost reductions in checking (lower C) reduce long-run p_t^{check} and termination hazards.
4. *Patience (discounting).* Greater effective patience δ (e.g., longer relationships, lower churn) reduces the total inspection frequency needed to deter (lower p^{check} at the sender margin) and, in the one-step benchmark, lowers λ^* (Section 8).
5. *Deterrence windows.* Publicly salient inspections create short-lived drops in deception (Assumption D): immediately after visible checks, measured deception rates fall before reverting (Lemma 5.3).

Proxies for B include stakes per interaction (claim size, shipment value, transaction margin). *Proxies for C* include staff/time costs, tooling, or automation adoption. *Patience* δ is proxied by expected relationship length or churn. Identification strategies: (i) event studies around inspection policy shocks or publicized sweeps; (ii) difference-in-differences when audit transparency or cost changes for a subset of units; (iii) instrument inspection costs via technology rollouts or staffing shifts. The model predicts stronger tapering and lower hazards where costs fall (lower C) or patience rises (higher δ).

If inspection outcomes are not publicly disclosed when truth is verified (silent audits), the public λ_t updates more slowly (Remark in Section 4), weakening tapering in p_t^{check} . If terminations are rare (punishment phase), equivalence (Proposition 6.3) implies similar cutoff behavior provided punishment is visible and finite.

Theorem 9.1. *At public state $x = (\lambda, \mu)$ under stationary policies (σ, ρ) , per-period expected surplus equals*

$$S(x) = B[1 - (1 - \lambda)\sigma(x)] - C p^{\text{check}}(x), \quad (9.1)$$

where $p^{\text{check}}(x) = \mu + (1 - \mu)\rho(x)$. Total welfare is

$$W(x_0) = \mathbb{E}_{x_0} \left[\sum_{t \geq 0} \delta^t S(X_t) \right],$$

and the termination hazard from (8.2) satisfies $h(x) = p^{\text{check}}(x)(1 - \lambda)\sigma(x)$. Along any stationary PBE path, W is (i) strictly increasing in δ , and (ii) weakly decreasing in C . The sign of $\partial W / \partial B$ is in general ambiguous.

Proof. In each period, surplus equals the receiver’s benefit when the realized action is truth minus the inspection cost. Under stationary policies, $\Pr(\text{truth at } x) = 1 - (1 - \lambda)\sigma(x)$ and the expected inspection cost is $C p^{\text{check}}(x)$, yielding (9.1). Discounted summation gives W . Monotone comparative statics in δ and C are immediate from linearity of $S(x)$ in these parameters and the fact that policies are well defined by equilibrium for each parameter value. The ambiguity in B follows because B raises the static gain on truthful periods but shifts equilibrium policies and thus h . $\square$

Corollary 9.2. *Fix a stationary PBE selection for each B . If the induced change in the hazard is small enough that*

$$\frac{\partial}{\partial B} \mathbb{E}[(1 - \Lambda_t)\sigma(X_t)] \leq \frac{1}{B} \mathbb{E}[1 - (1 - \Lambda_t)\sigma(X_t)] \quad \text{for all } t,$$

then $\partial W/\partial B \geq 0$. In particular, holding policies fixed (partial equilibrium) gives $\partial W/\partial B > 0$.

First, when the temptation-to-punishment ratio B/C is high, inspections are front-loaded: the inspection rate (and thus the termination hazard) is highest early on and then fades as the sender accumulates honesty reputation. Second, inspections taper as reputation improves: the receiver’s propensity to check declines with the public belief that the sender is honest, so the hazard falls mechanically along the relationship. Third, terminations bunch after inspection spurts: when the receiver temporarily raises inspections—because recent behavior or parameters make deception more attractive—the hazard spikes, producing a wave of exits, and then decays back as honesty beliefs recover and inspections ease.

10 Robustness and Extensions

We briefly discuss variants of the information/technology that preserve the paper’s qualitative conclusions and indicate how the key objects (belief recursion, indifference conditions, hazards) adjust.

10.1 Silent audits

This subsection summarizes the partial-disclosure environment in which truthful checks become public with probability $q \in [0, 1]$. The full Bayes maps, transition weights (disclosure vs. non-disclosure), and the modified indifference conditions are derived in the Online Appendix, which also proves continuity of stationary PBE in q and shows that $q \uparrow 1$ recovers the fully transparent benchmark while $q \downarrow 0$ approaches the no-public-signal limit.

We only record the comparative statics used in the main text: transparency tightens discipline. In particular, Proposition 10.1 implies that increasing q (i) lowers the receiver’s inspection cutoff in λ , (ii) shrinks the sender’s deception region, and (iii) weakly reduces the termination hazard $h(\lambda) = p^{\text{check}}(\lambda)(1 - \lambda)\sigma(\lambda)$ at each λ .

We study a partial-transparency device in which checks are conducted privately but a fraction of *truthful* checks are later disclosed to the public. Formally, when the receiver checks,

she privately learns the period outcome; if the sender was truthful, a disclosure occurs with probability $q \in (0, 1)$; if the sender deceived and is checked, the relationship terminates and that termination is publicly observed; trust generates no disclosure. Thus the only public objects are (i) disclosure vs. no-disclosure and (ii) termination; both players' actions remain privately observed.

This mechanism preserves a one-dimensional public state: the belief λ that the sender is the committed honest type. A disclosure raises λ ; no disclosure—conditional on survival—updates λ by Bayes because it can arise from trust, from a non-disclosed truthful check, or from unchecked deception. Hence λ follows a Markov recursion with two posteriors, $\lambda^1(\lambda)$ after disclosure and $\lambda^0(\lambda)$ after no disclosure.⁶ The public termination hazard at belief λ is the same as in the benchmark:

$$h(\lambda) = p^{\text{check}}(\lambda) (1 - \lambda) \sigma(\lambda),$$

since deception must be both present and checked to terminate.

Equilibrium analysis proceeds in this public state exactly as in the benchmark with public inspections. the Online Appendix proves: (i) *Existence*. A stationary PBE exists with a *receiver cutoff* in λ and a sender policy that is monotone in the effective discipline term generated by (a) public disclosure following truthful checks and (b) termination following checked deception. (ii) *Continuity in q* . Equilibria depend continuously on the disclosure probability: as $q \uparrow 1$ the model converges to the fully transparent benchmark in the main text (truthful checks always public), while as $q \downarrow 0$ public disclosures vanish and reputational discipline weakens toward the fully private-monitoring benchmark. (iii) *Comparative statics*. Increasing q strengthens the deterrence wedge in the sender's indifference, reduces the deception region (and the inspection intensity needed to sustain honesty), and lowers the long-run termination hazard along a given belief path.

Economically, even infrequent disclosures realign incentives by making continuation depend on publicly improved beliefs after disclosed truthful checks—exactly the transparency lever used in audits, certification schemes, and platform governance.

10.2 Noisy checks

We allow publicly observed checks to be imperfect. Conditional on a check, a binary signal $s \in \{\text{good}, \text{bad}\}$ is drawn with $\Pr(\text{good} \mid \text{truth}, \text{check}) = \pi_T$, $\Pr(\text{bad} \mid \text{deceive}, \text{check}) = \pi_D$, and precision $\kappa := \pi_T + \pi_D - 1 \in [0, 1]$ ($\kappa = 1$ recovers perfect verification; $\kappa = 0$ is uninformative). We keep the rule that a *bad* triggers termination (or the punishment phase) only if deception was present; *good* continues with Bayesian updating. The public state remains Markov, and the full Bayes maps/branch probabilities are in the Online Appendix.

Relative to the baseline, keep (5.1)–(5.2) but (i) replace the truthful-on-check term by

$$p^{\text{T|check}}(\lambda, \mu) \rightsquigarrow \pi_T \left[\lambda + (1 - \lambda) (1 - \sigma(\lambda, \mu)) \right],$$

⁶Explicit formulas in the Online Appendix.

and (ii) scale the detection/termination branch by π_D , so the per-period hazard becomes (Here $\kappa \equiv \pi_T + \pi_D - 1 \in [0, 1]$; the detection branch scales by π_D , so $h_\kappa(\lambda) = p^{\text{check}}(\lambda)(1 - \lambda)\sigma(\lambda)\pi_D$, consistent with Online Appendix OA2.)

$$h_\kappa(\lambda) = p^{\text{check}}(\lambda)(1 - \lambda)\sigma(\lambda)\pi_D, \quad (10.1)$$

with $h_\kappa \rightarrow h$ as $\kappa \uparrow 1$.⁷

Noisy verification and partial disclosure act as transparency devices. Let $q \in [0, 1]$ be the probability that a truthful check becomes public (silent audits; the Online Appendix). Both κ and q preserve a Markov public state and shift the same indifference system as a reduction in C or an increase in δ .

Proposition 10.1. *Fix (B, C, δ) . In any stationary PBE: (i) if $q_2 > q_1$, the receiver’s inspection cutoff in λ weakly decreases, the sender’s deception region shrinks, and the hazard $h(\lambda) = p^{\text{check}}(\lambda)(1 - \lambda)\sigma(\lambda)$ is weakly lower at each λ ; (ii) if $\kappa_2 > \kappa_1$, the inspection cutoff weakly decreases and $h_\kappa(\lambda)$ in (10.1) moves weakly toward the perfect-verification benchmark as $\kappa \uparrow 1$. Policies depend continuously on (q, κ) .*

Sketch. Differentiate the sender/receiver indifference conditions in q or κ : higher transparency increases the weight on publicly realized good news and/or detection probability, shifting Δ_r down in λ and Δ_s up in μ ; single-crossing then lowers the inspection cutoff and shrinks deception. Continuity follows from continuity of branch weights and the fixed-point map.

10.3 Fully private monitoring benchmark

When both players’ actions are privately observed and no audit outcome is ever made public, a low-dimensional public state does not exist: neither player observes how outsiders’ beliefs evolve, higher-order beliefs matter, and a stationary PBE generally cannot be expressed as a recursion in (λ, μ) . In particular, the only publicly observable event—termination after a detected deception—need not occur on path for long stretches, so Bayes’ rule provides no common update. This makes the fully private benchmark analytically fragile and non-recursive.

Two results from the Online Appendix clarify what survives and how to restore tractability. First, with *no* public information, stationary self-confirming equilibria (SCE) are easy to characterize but degenerate dynamically: the strategic sender strictly prefers deception in every period whenever $\delta < 1$, and the receiver’s inspection intensity is pinned down by the static inequality $R(1 - \theta) \geq C$. In the knife-edge case $R(1 - \theta) = C$ any inspection rate is consistent; otherwise the receiver either never inspects or always inspects, and the termination hazard is respectively 0 or front-loaded. This benchmark is useful as a robustness check, but it lacks the reputational cutoffs and tapering that are central in our main analysis.

Second, tractability (and meaningful reputation dynamics) is restored by injecting a *minimal* amount of public information. We show that an ε -frequency public “alarm” with

⁷Posteriors after a “good” signal are in the Online Appendix.

likelihood ratio $\kappa > 1$ favoring deception creates a one-dimensional public belief state λ and delivers a stationary PBE with the same cutoff logic as in the body of the paper. Moreover, equilibria depend continuously on ε and converge to a well-defined limit as $\varepsilon \downarrow 0$.

In Online Appendix we obtain the same conclusions when *truthful checks* are silently audited but disclosed with probability $q \in (0, 1)$: the public belief is again Markov, stationary PBE exist, and policies vary continuously in q up to the fully transparent benchmark $q \uparrow 1$. In both variants, the sender's indifference equates B to a discounted *discipline term* that scales with the public probability of either being checked or generating public news; this term is strictly positive even when the receiver chooses not to inspect, which is why infinitesimal public information suffices to reinstate reputational cutoffs.

For interpretation, the minimal-revelation devices correspond to practices such as randomized spot audits that publish outcomes, occasional whistleblower disclosures, or policy-driven transparency where a fraction of truthful inspections is publicized. Any such design that yields rare but informative public signals (either directly revealing truth or tilting alarms toward deception) maintains a Markov public state and preserves monotone best responses and cutoff comparative statics for inspection and deception. Thus, while the fully private benchmark highlights the limits of reputation without public signals, our main results apply verbatim once even a vanishing amount of public information is present, providing a clean rationale for transparency policies in auditing, certification, and platform moderation.

10.4 Discrete–continuous link

Consider a continuous-time (CT) version in which, at public honesty λ , the receiver inspects with intensity $r(\lambda)$ and the sender deceives with intensity $\sigma(\lambda)$; detected deception terminates. Let $\beta > 0$ be the CT discount rate, and let v_s, v_r solve the stationary HJB system associated with these intensities and the Bayes jump maps.

Discretize time with step $\Delta > 0$, discount $\delta_\Delta := e^{-\beta\Delta}$, and per-period probabilities

$$p_\Delta^{\text{check}}(\lambda) := r(\lambda)\Delta, \quad \sigma_\Delta(\lambda) := \sigma(\lambda)\Delta,$$

truncated to $[0, 1]$, with the same Bayes updates as in the baseline model. Let (V_s^Δ, V_r^Δ) be the discrete-time values under stationary policies $(\sigma_\Delta, \rho_\Delta)$ and let $(\lambda_\Delta^*, \mu_\Delta^*)$ denote the joint mixing cutoffs.

Lemma 10.2. *As $\Delta \downarrow 0$, the discrete-time Bellman operators converge to the CT resolvent, and the values converge uniformly:*

$$\lim_{\Delta \downarrow 0} \|V_i^\Delta - v_i\|_\infty = 0, \quad i \in \{s, r\}.$$

Moreover, every sequence of stationary PBE selections admits a subsequence for which the cutoffs and hazards converge:

$$(\lambda_\Delta^*, \mu_\Delta^*) \rightarrow (\lambda^*, \mu^*), \quad h_\Delta(\lambda) = p_\Delta^{\text{check}}(\lambda)(1 - \lambda)\sigma_\Delta(\lambda) \Rightarrow h(\lambda) = r(\lambda)(1 - \lambda)\sigma(\lambda),$$

with (λ^, μ^*) a stationary CT cutoff (MPE) and h the CT termination intensity.*

Proof sketch. The discrete transition kernel is an Euler discretization of the CT generator; the Bellman operator converges to the CT resolvent as $\Delta \rightarrow 0$. Uniform convergence of values follows from standard resolvent–HJB arguments under bounded intensities. Stationary cutoffs solve a continuous system of indifference conditions in the values; by continuity and compactness, subsequential convergence holds, and the discrete hazard h_Δ converges pointwise to h . $\square$

11 Conclusion

This paper analyzes a repeated sender–receiver environment in which inspection decisions are public while the sender’s actions are private unless checked. The public observability of inspections yields a low-dimensional belief state (honesty and vigilance) on which we construct a stationary PBE and prove existence (Lemma 7.1). Equilibrium policies take a cutoff form: the receiver checks when honesty beliefs are sufficiently low and the sender deceives only when vigilance beliefs are sufficiently low, with monotone best responses (Lemma 5.2) and locally unique thresholds pinned down by two indifference equations. In a benchmark with short deterrence windows we obtain closed-form expressions for the total inspection rate required to deter and for the honesty cutoff (Lemma 5.3). We further show that finite, publicly observed punishments implement the same cutoffs as immediate termination (Proposition 6.3), so discipline is sustained by reputational cutoffs rather than by absorption per se.

The model delivers transparent comparative statics and hazard predictions. Greater temptation raises the checking region while higher inspection costs shrink it; greater patience reduces the inspection intensity needed to sustain honesty (Proposition 8.1). Inspection intensity tapers with accumulated honesty, and terminations bunch following inspection spurts (Corollary 8.2). These results map directly to auditing, certification, platform moderation, and insurance contexts where inspections are visible but violations are revealed only upon inspection. For policy and design, the analysis suggests front-loaded inspections when B/C is high, a gradual taper as honesty accumulates, and the possibility of replacing one-strike termination with finite, visible probation without altering equilibrium cutoffs—potentially lowering administrative costs while preserving deterrence.

Several limitations point to fruitful extensions. We show robustness to silent audits and noisy checks, yet fully private monitoring remains analytically intractable without minimal public signals; characterizing optimal information design that injects such signals is a natural next step. Endogenizing enforcement instruments (e.g., inspection cost via technology adoption) and punishment menus, allowing heterogeneous types or networked interactions (platform ecosystems), or moving to continuous time with time-varying inspection rates could sharpen welfare and implementation results. On the empirical side, the model’s cutoffs and hazard predictions lend themselves to event studies and cost-shock designs; estimating the implied (B, C, δ) from audit and moderation data would help quantify the efficiency gains from reputational enforcement and the substitution between termination and finite probation.

A Proofs

Proof of Lemma 5.2. By Proposition 5.1(a)–(b), the continuation values V_s, V_r are monotone in (λ, μ) . Using (A.1)–(A.2), this implies Δ_r is decreasing in λ and Δ_s in μ , hence the stated policy monotonicities; the cutoff forms follow because binary best replies to single-crossing differences admit thresholds.

Fix a public state $x = (\lambda, \mu)$ and write

$$\begin{aligned} p^{\text{check}}(x) &:= \mu + (1 - \mu) \rho(x), \\ p^{\text{T|check}}(x) &:= \lambda + (1 - \lambda) (1 - \sigma(x)). \end{aligned}$$

From (4.1)–(4.2), the action-contingent values at x are

$$\begin{aligned} V_s^{\text{deceive}}(x) &= B + \delta (1 - p^{\text{check}}(x)) V_s(\lambda, 0), \\ V_s^{\text{truth}}(x) &= \delta [p^{\text{check}}(x) V_s(\lambda^+(x), \mu^+(x)) + (1 - p^{\text{check}}(x)) V_s(\lambda, 0)], \\ V_r^{\text{check}}(x) &= B p^{\text{T|check}}(x) - C + \delta p^{\text{T|check}}(x) V_r(\lambda^+(x), \mu^+(x)), \\ V_r^{\text{trust}}(x) &= B [1 - (1 - \lambda) \sigma(x)] + \delta V_r(\lambda, 0). \end{aligned}$$

Define the action-difference functions

$$\begin{aligned} \Delta_r(x) &:= V_r^{\text{check}}(x) - V_r^{\text{trust}}(x), \\ \Delta_s(x) &:= V_s^{\text{deceive}}(x) - V_s^{\text{truth}}(x). \end{aligned}$$

A direct subtraction yields

$$\Delta_r(x) = R(1 - \lambda) \sigma(x) - C + \delta [V_r(\lambda^+(x), \mu^+(x)) - V_r(\lambda, 0)], \quad (\text{A.1})$$

and

$$\Delta_s(x) = B - \delta p^{\text{check}}(x) [V_s(\lambda^+(x), \mu^+(x)) - V_s(\lambda, 0)]. \quad (\text{A.2})$$

Step 1. Fix μ and $\lambda' < \lambda''$. Compare $\Delta_r(\lambda'', \mu)$ vs. $\Delta_r(\lambda', \mu)$ using (A.1).

- (i) The current detection term $R(1 - \lambda) \sigma(\lambda, \mu)$ is weakly decreasing in λ , since $1 - \lambda$ falls in λ and $\sigma \in [0, 1]$.
- (ii) For the continuation term, note that the Bayes map $\lambda \mapsto \lambda^+(\lambda, \mu) = \frac{\lambda}{\lambda + (1 - \lambda)(1 - \sigma(\lambda, \mu))}$ is increasing and concave in λ (fractional linear), hence the increment $\lambda^+(\lambda, \mu) - \lambda$ is weakly decreasing in λ .

By hypothesis, V_r is increasing in λ and in μ , while $\mu^+(\lambda, \mu) = \frac{\mu}{\mu + (1 - \mu) \rho(\lambda, \mu)} \in [\mu, 1]$ and the trust-posterior on the receiver equals 0. Therefore

$$V_r(\lambda^+(\lambda'', \mu), \mu^+(\lambda'', \mu)) - V_r(\lambda'', 0) \leq V_r(\lambda^+(\lambda', \mu), \mu^+(\lambda', \mu)) - V_r(\lambda', 0),$$

i.e., the continuation gain from a truthful public check is weakly smaller at higher λ . Combining (i)–(ii) in (A.1) yields $\Delta_r(\lambda'', \mu) \leq \Delta_r(\lambda', \mu)$, so $\Delta_r(\cdot, \mu)$ is weakly decreasing in λ .

Since the receiver's best reply is $\rho(\lambda, \mu) \in \arg \max\{\Delta_r(\lambda, \mu), 0\}$ for each (λ, μ) , it follows that $\rho(\lambda, \mu)$ is weakly decreasing in λ .

Step 2. Fix λ and $\mu' < \mu''$. Consider (A.2). Write

$$\Theta(\mu) := p^{\text{check}}(\lambda, \mu) \left[V_s(\lambda^+(\lambda, \mu), \mu^+(\lambda, \mu)) - V_s(\lambda, 0) \right].$$

We show $\Theta(\mu)$ is weakly increasing in μ , which implies $\Delta_s(\lambda, \mu) = B - \delta \Theta(\mu)$ is weakly decreasing in μ . First, $p^{\text{check}}(\lambda, \mu) = \mu + (1 - \mu) \rho(\lambda, \mu)$ is weakly increasing in μ (holding ρ at the respective states). Second, $\mu^+(\lambda, \mu) = \frac{\mu}{\mu + (1 - \mu) \rho(\lambda, \mu)}$ is weakly increasing in μ , and by hypothesis V_s is weakly decreasing in μ and increasing in λ . Therefore

$$V_s(\lambda^+(\lambda, \mu''), \mu^+(\lambda, \mu'')) - V_s(\lambda, 0) \leq V_s(\lambda^+(\lambda, \mu'), \mu^+(\lambda, \mu')) - V_s(\lambda, 0),$$

so the sender's continuation *loss* from a truthful public check weakens as μ rises. Multiplying a weakly larger p^{check} by a weakly smaller nonnegative bracket preserves weak increase of the product (values are bounded), hence $\Theta(\mu'') \geq \Theta(\mu')$ and thus $\Delta_s(\lambda, \mu'') \leq \Delta_s(\lambda, \mu')$. The sender's best reply is $\sigma(\lambda, \mu) \in \arg \max\{\Delta_s(\lambda, \mu), 0\}$, hence $\sigma(\lambda, \mu)$ is weakly decreasing in μ .

Step 3. For each fixed μ , $\Delta_r(\cdot, \mu)$ is weakly decreasing, so the receiver's binary choice admits a threshold: there exists a (possibly set-valued) cutoff $\lambda^*(\mu)$ such that the receiver checks iff $\lambda \leq \lambda^*(\mu)$ (mixing when equality holds). Likewise, for each fixed λ , $\Delta_s(\lambda, \cdot)$ is weakly decreasing, so there exists $\mu^*(\lambda)$ such that the sender deceives iff $\mu \leq \mu^*(\lambda)$ (mixing at equality).

This proves the lemma. $\square$

Proof of Lemma 5.3. Let $x^* = (\lambda^*, \mu^*)$ be a sender-mixing public state in the termination benchmark, so that the Bellman equalities at x^* hold and both actions are used with positive probability. Recall the sender's and receiver's indifference equations (cf. (5.1)–(5.2)):

$$B = \delta p^{\text{check}}(x^*) \left[V_s(\lambda^+(x^*), \mu^+(x^*)) - V_s(\lambda^*, 0) \right], \quad (\text{A.3})$$

$$C = \delta \left[p^{\text{Tcheck}}(x^*) V_r(\lambda^+(x^*), \mu^+(x^*)) - V_r(\lambda^*, 0) \right]. \quad (\text{A.4})$$

Step 1. By Assumption D(i) (local stationarity at the sender's cutoff), the *continuation wedge* created by a truthful, publicly observed check equals a perpetuity worth $B/(1 - \delta)$:

$$V_s(\lambda^+(x^*), \mu^+(x^*)) - V_s(\lambda^*, 0) = \frac{B}{1 - \delta}. \quad (\text{A.5})$$

Substituting (A.5) into the sender indifference (A.3) gives

$$B = \delta p^{\text{check}}(x^*) \frac{B}{1 - \delta},$$

hence

$$p^{\text{check}*} := p^{\text{check}}(x^*) = \frac{1 - \delta}{\delta},$$

which is (5.3). Noting that $p^{\text{check}}(\lambda, \mu) = \mu + (1 - \mu) \rho(\lambda, \mu)$ yields the stated equivalent form $\mu + (1 - \mu) \bar{\rho} = (1 - \delta)/\delta$ at the sender's cutoff.

Step 2. Suppose that a truthful, publicly observed check at x^* *reduces next-period deception* by the strategic sender from $\bar{\sigma}$ to $\bar{\sigma}'$, with all other components of the continuation environment locally unchanged (Assumption D(ii), “one-step deterrence”). Then the receiver's *one-period* gain in expected benefit at $t+1$ equals

$$B \cdot \underbrace{(1 - \lambda^*)}_{\text{prob. strategic sender}} \cdot \underbrace{(\bar{\sigma} - \bar{\sigma}')}_{\text{drop in deception}},$$

and stationarity propagates this increment forward. Using the sender's wedge identity from Step 1, the implied *continuation* gain can be written in closed form as

$$V_r(\lambda^+(x^*), \mu^+(x^*)) - V_r(\lambda^*, 0) = \frac{p^{\text{check}*}}{p^{\text{T|check}}(x^*)} B (1 - \lambda^*) (\bar{\sigma} - \bar{\sigma}'), \quad (\text{A.6})$$

i.e., the Bayes branch weight $p^{\text{T|check}}(x^*)$ cancels when we translate a one-period improvement into a continuation increment under the equilibrium survival/termination mixture.⁸

Substituting (A.6) in the receiver indifference (A.4) yields

$$C = \delta p^{\text{T|check}}(x^*) \cdot \frac{p^{\text{check}*}}{p^{\text{T|check}}(x^*)} B (1 - \lambda^*) (\bar{\sigma} - \bar{\sigma}') = \delta B p^{\text{check}*} (1 - \lambda^*) (\bar{\sigma} - \bar{\sigma}').$$

Solving for λ^* gives

$$\lambda^* = 1 - \frac{C}{\delta B p^{\text{check}*} (\bar{\sigma} - \bar{\sigma}')},$$

which is (5.4). The comparative statics in the last sentence follow immediately: λ^* is decreasing in the deterrence gap $(\bar{\sigma} - \bar{\sigma}')$ and increasing in C/B . The formula is meaningful whenever the right-hand side lies in $(0, 1)$. $\square$

OA1 Silent audits

We study the case in which the receiver's inspection decision is publicly observed, but the *outcome* of a truthful check is not disclosed: only detected deception (**deceive, check**) becomes public via termination (absorbing). Thus the public state remains two-dimensional $x = (\lambda, \mu)$: honesty belief λ and vigilance belief μ . The key change relative to the main benchmark is that on truthful-check branches, the public honesty belief does *not* jump: $\lambda_{\text{pub}}^+ = \lambda$ (while the receiver's *private* posterior increases).

⁸Formally, write the receiver's Bellman at x^* under the two continuations (truthful-check vs. no-check) and isolate the part that varies with the next-period deception rate. The mapping from a one-period change in deception probability to a continuation change is linear; the factor $p^{\text{check}*}/p^{\text{T|check}}(x^*)$ is obtained by substituting the sender's indifference (Step 1) to eliminate the common survival multiplier. This is the same “equalization” device that delivers (5.3).

At $x = (\lambda, \mu)$ the total check probability is $p^{\text{check}}(x) = \mu + (1 - \mu) \rho(x)$. When a check occurs and the sender is truthful, the public next state is $(\lambda, \mu^+(x))$; when a check occurs and the sender deceives, termination occurs; with no check the next state is $(\lambda, 0)$ (vigilance belief resets because no public check happened). Formally, replace every occurrence of $V_s(\lambda^+(x), \cdot)$ on truthful-check branches in the public Bellman system by $V_s(\lambda, \cdot)$; the vigilance posterior $\mu^+(x)$ remains as in the main model because checks are public.

Let $V_s(x)$ and $V_r(x)$ denote continuation values. Then

$$V_s(\lambda, \mu) = \max_{\sigma \in [0,1]} \left\{ \sigma B + \delta \left[(1 - p^{\text{check}}(x)) V_s(\lambda, 0) + \underbrace{p^{\text{check}}(x) (1 - \sigma) V_s(\lambda, \mu^+(x))}_{\text{truthful check: no public honesty jump}} \right] \right\}, \quad (\text{OA1.1})$$

$$V_r(\lambda, \mu) = \max_{\rho \in [0,1]} \left\{ \rho \left(B p^{\text{T|check}}(x) - C + \delta p^{\text{T|check}}(x) \underbrace{V_r(\lambda, \mu^+(x))}_{\text{truthful check continues}} \right) + (1 - \rho) \left(B [1 - (1 - \lambda) \sigma(x)] + \delta V_r(\lambda, 0) \right) \right\}. \quad (\text{OA1.2})$$

Here $p^{\text{T|check}}(x) = \lambda + (1 - \lambda) [1 - \sigma(x)]$ is the probability a check verifies truth.

At any belief x^* where both players mix, sender and receiver indifferences are

$$B = \delta p^{\text{check}}(x^*) \left[V_s(\lambda^*, \mu^+(x^*)) - V_s(\lambda^*, 0) \right], \quad (\text{OA1.3})$$

$$C = \delta p^{\text{T|check}}(x^*) \left[V_r(\lambda^*, \mu^+(x^*)) - V_r(\lambda^*, 0) \right]. \quad (\text{OA1.4})$$

Compared to full disclosure, the sender's "truthful-check continuation gain" replaces

$$V_s(\lambda^+(x^*), \mu^+(x^*)) - V_s(\lambda^*, 0)$$

by $V_s(\lambda^*, \mu^+(x^*)) - V_s(\lambda^*, 0)$, which is weakly smaller because V_s is (weakly) increasing in honesty belief.

Theorem OA1.1. *Under bounded payoffs and $\delta \in (0, 1)$, the silent-audit model admits a stationary PBE in the public state $x = (\lambda, \mu)$ with (possibly mixed) policies (σ, ρ) and bounded values (V_s, V_r) solving (OA1.1)–(OA1.2). At any mixed x^* , (OA1.3)–(OA1.4) hold.*

Proof. Fix measurable (σ, ρ) . The right-hand sides of (OA1.1)–(OA1.2) define a contraction mapping on the product of bounded functions over $[0, 1]^2$ with modulus δ (dependence on (V_s, V_r) is affine and multiplied by δ), yielding unique bounded continuous values (V_s, V_r) . Pointwise best-reply correspondences are nonempty, convex-valued, and u.h.c. by Berge's theorem. On the compact convex space of measurable policies $[0, 1]^{[0,1]^2} \times [0, 1]^{[0,1]^2}$ (product topology), Kakutani–Fan–Glicksberg gives a fixed point; measurable selections exist by Kuratowski–Ryll–Nardzewski. Bayes consistency is built into μ^+ and the termination branch. $\square$

Lemma OA1.2. *If V_s and V_r are (weakly) increasing in λ and μ , then the receiver's best reply is (weakly) decreasing in λ and (weakly) increasing in μ ; the sender's best reply is (weakly) decreasing in μ and increasing in λ only through $p^{\text{T|check}}$ and p^{check} . Hence stationary PBE admit cutoff characterizations in λ for fixed μ , and in μ for fixed λ .*

Proof. Compare the receiver's "check – trust" difference implied by (OA1.2):

$$\Delta_r(x) = -C + \delta p^{\text{T|check}}(x) \left[V_r(\lambda, \mu^+(x)) - V_r(\lambda, 0) \right].$$

This is decreasing in λ (since $p^{\text{T|check}}$ is increasing in λ but the bracketed continuation gain does not depend on λ under silent audits) and increasing in μ via $\mu^+(x)$. For the sender, the difference "truth – deceive" from (OA1.1) is

$$\Delta_s(x) = -B + \delta p^{\text{check}}(x) \left[V_s(\lambda, \mu^+(x)) - V_s(\lambda, 0) \right],$$

which is increasing in μ via p^{check} and $\mu^+(x)$; dependence on λ only enters through $p^{\text{check}}, p^{\text{T|check}}$. Monotonicity yields the stated cutoffs. $\square$

Proposition OA1.3. *Let $\lambda_{\text{full}}^*(\mu)$ be the receiver's cutoff in the full-disclosure model and $\lambda_{\text{silent}}^*(\mu)$ the cutoff under silent audits. Then for all $\mu \in [0, 1]$,*

$$\lambda_{\text{silent}}^*(\mu) \geq \lambda_{\text{full}}^*(\mu),$$

with strict inequality whenever the truthful-check continuation gain $V_s(\lambda^+(x), \mu^+(x)) - V_s(\lambda, 0)$ is strictly increasing in λ at $x = (\lambda_{\text{full}}^(\mu), \mu)$.*

Proof. Fix μ and evaluate the sender's margin at $\lambda = \lambda_{\text{full}}^*(\mu)$ and the policies from the full-disclosure equilibrium. In that model,

$$B = \delta p^{\text{check}}(x) \left[V_s(\lambda^+(x), \mu^+(x)) - V_s(\lambda, 0) \right].$$

Under the silent-audit recursion, the same (σ, ρ) and λ yield the *weakly smaller* right-hand side

$$\delta p^{\text{check}}(x) \left[V_s(\lambda, \mu^+(x)) - V_s(\lambda, 0) \right],$$

because $\lambda^+(x) \geq \lambda$ and V_s is increasing in λ (Blackwell improvement of the transition). Hence, holding λ fixed, the sender would strictly prefer deception under silent audits unless p^{check} increases. Since the receiver's best reply is decreasing in λ (Lemma OA1.2), raising p^{check} at a mixing point requires *weakly* higher ρ , which in a cutoff policy corresponds to a *weakly higher* cutoff λ^* . The strict case follows when the truthful-check continuation gain is strictly larger under full disclosure at that λ . $\square$

Let a truthful check be publicly disclosed with probability $q \in [0, 1]$ (and otherwise silent). Then the truthful-check continuation term in (OA1.1) is replaced by

$$q V_s(\lambda^+(x), \mu^+(x)) + (1 - q) V_s(\lambda, \mu^+(x)),$$

and analogously in the receiver's Bellman. Existence and monotone BR follow as above. Moreover:

Proposition OA1.4. *Let $\lambda^*(q; \mu)$ be the receiver’s cutoff at transparency level $q \in [0, 1]$. Then $\lambda^*(q; \mu)$ is weakly decreasing in q for every μ . Consequently, the equilibrium termination hazard is weakly decreasing in q for fixed primitives (B, C, δ) .*

Proof. The Bellman operator with parameter q' Blackwell–dominates that with q whenever $q' > q$, because the truthful–check continuation is a convex combination that increases with q (keeping μ –transitions fixed). Hence the fixed–point values (V_s, V_r) are weakly increasing in q (parametric contraction + monotone operator). In the sender’s indifference (OA1.3), a larger q raises the right–hand side at fixed $(\lambda, \mu, \sigma, \rho)$; thus to maintain equality with constant B , a *smaller* p^{check} suffices. Since p^{check} is decreasing in the receiver’s cutoff (Lemma OA1.2), λ^* must (weakly) fall with q . The hazard $h(x) = p^{\text{check}}(x) (1 - \lambda) \sigma(x)$ then (weakly) falls with q along the equilibrium policy. $\square$

Relative to full public disclosure, silent audits weaken the public informativeness of truthful checks and therefore expand the checking region (higher λ^*) and (weakly) raise hazards. Allowing partial transparency $q \in (0, 1)$ interpolates linearly in the truthful–check term and yields a clean monotone comparative static: more transparency $\Rightarrow$ less checking needed to deter, lower hazards, and (weakly) higher values.

OA2 Noisy checks

We extend the benchmark with publicly observed *checks* by allowing the check to return a binary *signal* $s \in \{\text{good}, \text{bad}\}$ that is only *informative* about the sender’s period action. Conditional on a check, the signal technology is

$$\begin{aligned}\Pr(\text{good} \mid \text{truth}, \text{check}) &= \pi_T, \\ \Pr(\text{bad} \mid \text{deceive}, \text{check}) &= \pi_D, \\ \kappa &:= \pi_T + \pi_D - 1 \in [0, 1],\end{aligned}$$

where κ is precision ($\kappa = 1$ is perfect verification; $\kappa = 0$ is uninformative). The action *check* remains publicly observed. As before, (*deceive*, *check*) may trigger termination (benchmark) or a finite punishment phase (extension).

We analyze two natural disclosure regimes:

Regime A (deception–only termination). A *bad* signal leads to termination *only if deception was present* that period; equivalently, bad signals on truthful periods are automatically overturned (e.g., via review) and do not terminate. This matches the replacement rules stated in the main text.

Regime B (terminate on any bad). Any *bad* signal triggers termination, even if the sender was truthful (false positives cause erroneous termination).

Unless otherwise noted, we present formulas under Regime A (to mirror the main-text mapping) and then state the simple modifications for Regime B.

OA2.1 Belief updates and hazards

Let λ be the public belief that the sender is the committed honest type and let $\sigma(x) \in [0, 1]$ be the deception probability of a strategic sender at public state x (which can be λ alone or (λ, μ) depending on your baseline). When a *check* occurs, the probability that the period action was *truth* equals

$$\Pr(\text{truth} \mid \text{check}, x) = \lambda + (1 - \lambda) [1 - \sigma(x)]. \quad (\text{OA2.1})$$

A *good* signal can arise from truth (probability π_T) and, if deception occurred, from a false negative (probability $1 - \pi_D$). Bayes' rule therefore yields the posterior after a *good* signal:

$$\lambda_{\text{good}}^+(x) = \frac{\pi_T [\lambda + (1 - \lambda)(1 - \sigma(x))]}{\pi_T [\lambda + (1 - \lambda)(1 - \sigma(x))] + (1 - \pi_D)(1 - \lambda)\sigma(x)}. \quad (\text{OA2.2})$$

If *trust* is publicly observed (no check), there is no information about the sender's action that period, so the public honesty belief does not update: $\lambda^{\text{ns}} = \lambda$.⁹

The per-period *termination hazard* differs across regimes:

$$\text{Regime A:} \quad h(x) = p^{\text{check}}(x) [(1 - \lambda)\sigma(x)\pi_D], \quad (\text{OA2.3})$$

$$\text{Regime B:} \quad h(x) = p^{\text{check}}(x) [\lambda(1 - \pi_T) + (1 - \lambda)\sigma(x)\pi_D]. \quad (\text{OA2.4})$$

where $p^{\text{check}}(x) = \mu + (1 - \mu)\rho(x)$ is the public probability of a check against a strategic sender (as in the main text).

OA2.2 Bellman equations and indifference

Write $V_s(x)$ and $V_r(x)$ for continuation values at the public state x . If a *check* occurs and the signal is *good*, the relationship continues and next period's public state is $(\lambda_{\text{good}}^+(x), \mu^+)$; if *bad* and Regime A holds, termination reveals deception and stops. If *trust*, there is no signal and the state is $(\lambda, \mu^{\text{ns}})$.

Under Regime A, the sender's Bellman equation is

$$V_s(x) = \max \left\{ \overbrace{\delta \left[(1 - p^{\text{check}}(x)) V_s(\lambda, \mu^{\text{ns}}) + p^{\text{check}}(x) V_s(\lambda_{\text{good}}^+(x), \mu^+) \right]}^{\text{truth}}, \right. \\ \left. \underbrace{B + \delta \left[(1 - p^{\text{check}}(x)) V_s(\lambda, \mu^{\text{ns}}) + p^{\text{check}}(x) (1 - \pi_D) V_s(\lambda_{\text{good}}^+(x), \mu^+) \right]}_{\text{deceive}} \right\}. \quad (\text{OA2.5})$$

⁹If you track the vigilance reputation μ of the receiver (committed always-check type), it updates upon observing a check (via $\mu^+ = \frac{\mu}{\mu + (1 - \mu)\rho(x)}$) and collapses to zero upon observing trust.

The receiver's Bellman equation is

$$V_r(x) = \max \left\{ \overbrace{B \left[1 - (1 - \lambda)\sigma(x) \right] + \delta V_r(\lambda, \mu^{\text{ns}})}^{\text{trust}}, \right. \\ \left. \overbrace{B \left[\lambda + (1 - \lambda)(1 - \sigma(x)) \right] + R(1 - \lambda)\sigma(x)\pi_D - C + \delta V_r(\lambda_{\text{good}}^+(x), \mu^+)}^{\text{check}} \right\}. \quad (\text{OA2.6})$$

At any mixing state x^* (where both pure actions are used with positive probability), indifference yields

$$\text{Sender: } B = \delta p^{\text{check}}(x^*) \pi_D \left[V_s(\lambda_{\text{good}}^+(x^*), \mu^+) - V_s(\lambda, \mu^{\text{ns}}) \right], \quad (\text{OA2.7})$$

$$\text{Receiver: } C = R(1 - \lambda^*)\sigma(x^*)\pi_D + \delta \left(V_r(\lambda_{\text{good}}^+(x^*), \mu^+) - V_r(\lambda, \mu^{\text{ns}}) \right). \quad (\text{OA2.8})$$

The B -terms cancel in (OA2.8) because the expected stage benefit from truth is identical under trust and check.

If any *bad* triggers termination (even under truth), then in (OA2.5) replace the truth continuation by $\delta [(1 - p^{\text{check}}(x))V_s(\lambda, \mu^{\text{ns}}) + p^{\text{check}}(x)\pi_T V_s(\lambda_{\text{good}}^+(x), \mu^+)]$ and in (OA2.6) replace the check continuation analogously. The indifference conditions become

$$B = \delta p^{\text{check}}(x^*) \left[\pi_D V_s(\lambda_{\text{good}}^+(x^*), \mu^+) - (1 - \pi_T) V_s(\lambda, \mu^{\text{ns}}) \right], \quad (\text{OA2.9})$$

$$C = R \left[\lambda^*(1 - \pi_T) + (1 - \lambda^*)\sigma(x^*)\pi_D \right] + \delta \left(\pi_T V_r(\lambda_{\text{good}}^+(x^*), \mu^+) - V_r(\lambda, \mu^{\text{ns}}) \right). \quad (\text{OA2.10})$$

OA2.3 Existence and structure

Existence of a stationary PBE follows as in the benchmark: the policy-evaluation operator remains a δ -contraction (the only change is the continuation branch weights), and pointwise best replies are affine in own control; Berge + Kakutani-Fan-Glicksberg deliver a fixed point (details omitted to avoid repetition). Monotone best responses and cutoff inspection in the public honesty belief carry over provided V_r is increasing in λ (value monotonicity) and $p^{\text{check}}(\cdot)$ is weakly decreasing in λ .

Lemma OA2.1. *Fix $\pi_T, \pi_D \in [0, 1]$ with $\kappa = \pi_T + \pi_D - 1 \in [0, 1]$. Then $\lambda_{\text{good}}^+(x)$ in (OA2.2) is (i) strictly increasing in λ , (ii) weakly decreasing in $\sigma(x)$, (iii) weakly increasing in π_T , and (iv) weakly increasing in π_D . If $\kappa = 1$ (perfect checks), $\lambda_{\text{good}}^+ = \frac{\lambda}{\lambda + (1 - \lambda)[1 - \sigma(x)]}$, matching the perfect-verification benchmark.*

Proof. Each claim follows by direct differentiation of the rational expression in (OA2.2) (denominator positive) and the fact that a higher π_T or π_D improves the likelihood ratio of good in favor of truth. $\square$

Proposition OA2.2. *Under Regime A, as κ falls (i.e., π_T or π_D falls), the sender’s discipline term on the right-hand side of (OA2.7) weakens and the receiver’s informational gain in (OA2.8) weakens. Consequently, any inspection cutoff λ^* (if unique) shifts outward (higher), and the termination hazard $h(x)$ weakly falls pointwise. The same qualitative conclusions hold under Regime B.*

If bad triggers a finite, publicly visible punishment phase rather than termination, replace the termination branch in (OA2.5)–(OA2.6) with the corresponding punishment values $V_s^{\text{pun}}, V_r^{\text{pun}}$. The outcome-equivalence result from the main text carries over: for any termination equilibrium, there exists a finite punishment system that implements the same mixing cutoffs (the proof is identical with the noisy weights inserted).

(i) The main text’s replacement rules correspond to Regime A; if you prefer the more stringent Regime B, apply (OA2.4)–(OA2.10). (ii) In the limit $\kappa \uparrow 1$ we recover the perfect-verification model; in the opposite limit $\kappa \downarrow 0$ the signal is uninformative, the good posterior collapses to λ , and reputational discipline reverts toward that with silent audits of vanishing disclosure.

OA3 Private monitoring

This appendix studies a fully private-monitoring environment in which both actions (the sender’s *truth/deceive* and the receiver’s *trust/check*) are privately observed, but there is a small, exogenous *public* signal each period. The public signal is informative about the period’s outcome and preserves a low-dimensional public belief about the sender’s honesty; equilibrium strategies are Markov in this public belief.

Stage payoffs match the main text except we allow a contemporaneous compensation to the receiver upon detected deception: if (*deceive, check*) occurs, the receiver obtains $R - C$ and the relationship terminates; if (*truth, check*) occurs, she obtains $B - C$; if (*truth, trust*) occurs, she obtains B ; and if (*deceive, trust*) occurs, she obtains 0. The sender receives B when he deceives (regardless of inspection) and 0 when truthful.

OA3.1 Minimal Public Revelation

Time is discrete; discount factor $\delta \in (0, 1)$. Let λ_t denote the *public* belief that the sender is the committed honest type at the start of period t . Each period yields a binary public signal $S_t \in \{0, 1\}$ with small frequency $\varepsilon \in (0, 1)$ and likelihood ratio $\kappa > 1$ favoring deception:

$$\Pr(S_t = 1 \mid \text{truth}) = \varepsilon, \quad \Pr(S_t = 1 \mid \text{deceive}) = \kappa \varepsilon \quad (\leq 1).$$

Receiver checks and sender actions are private unless (*deceive, check*) occurs, which terminates the relationship and is publicly observed.

Let $\sigma(\lambda) \in [0, 1]$ be the strategic sender’s deception probability at public belief λ and $\rho(\lambda) \in [0, 1]$ the strategic receiver’s inspection probability (both depend on λ only). Because checks are private, the public state is one-dimensional (no public μ recursion). For notational

convenience, let $\bar{p}^{\text{check}}(\lambda) := \mu_0 + (1 - \mu_0) \rho(\lambda)$ denote the ex-ante check probability against a strategic sender (where $\mu_0 \in [0, 1]$ is the prior weight on a committed vigilant receiver).

Given mixing $\sigma(\lambda)$ at belief λ , Bayes' rule delivers the next public belief after the public signal:

$$\lambda^1(\lambda) \equiv \Pr(\text{honest} \mid S = 1, \lambda) = \frac{\lambda}{\lambda + (1 - \lambda) [\kappa \sigma(\lambda) + (1 - \sigma(\lambda))]}, \quad (\text{OA3.1})$$

$$\lambda^0(\lambda) \equiv \Pr(\text{honest} \mid S = 0, \lambda) = \frac{(1 - \varepsilon) \lambda}{(1 - \varepsilon) \lambda + (1 - \lambda) [1 - \kappa \varepsilon \sigma(\lambda) - \varepsilon (1 - \sigma(\lambda))]} \quad (\text{OA3.2})$$

Define the variant-specific continuation operators (used below) as the expected next-value under truth or deception conditional on survival:

$$\begin{aligned} \mathcal{G}_i^{\text{truth}}(\lambda) &:= \varepsilon V_i(\lambda^1(\lambda)) + (1 - \varepsilon) V_i(\lambda^0(\lambda)), \\ \mathcal{G}_i^{\text{deceive}}(\lambda) &:= \kappa \varepsilon V_i(\lambda^1(\lambda)) + (1 - \kappa \varepsilon) V_i(\lambda^0(\lambda)), \end{aligned}$$

for $i \in \{s, r\}$.

Let $V_s(\lambda)$ and $V_r(\lambda)$ be the (public-state) continuation values when the relationship is active at public belief λ . Checking is private; the only publicly visible resolution is termination upon detected deception.

If the sender is strategic and chooses **deceive** at λ , his value is

$$B + \delta (1 - \bar{p}^{\text{check}}(\lambda)) \mathcal{G}_s^{\text{deceive}}(\lambda),$$

where survival requires “no check” against a strategic sender. If he chooses **truth**, his value is

$$0 + \delta \mathcal{G}_s^{\text{truth}}(\lambda).$$

At any mixing belief λ^* ,

$$B = \delta \left[\mathcal{G}_s^{\text{truth}}(\lambda^*) - (1 - \bar{p}^{\text{check}}(\lambda^*)) \mathcal{G}_s^{\text{deceive}}(\lambda^*) \right]. \quad (\text{OA3.3})$$

If the receiver *checks*, her current payoff is

$$\underbrace{B \cdot [\lambda + (1 - \lambda)(1 - \sigma(\lambda))]}_{\text{truth branch}} + \underbrace{R \cdot (1 - \lambda)\sigma(\lambda)}_{\text{compensation on detected deception}} - C,$$

and her continuation equals $\delta \cdot [\lambda + (1 - \lambda)(1 - \sigma(\lambda))] \cdot \mathcal{G}_r^{\text{truth}}(\lambda)$ (since deception triggers termination). If she *trusts*, her current payoff is $B \cdot [\lambda + (1 - \lambda)(1 - \sigma(\lambda))]$ and continuation is $\delta \cdot (\lambda \cdot \mathcal{G}_r^{\text{truth}}(\lambda) + (1 - \lambda)\sigma(\lambda) \cdot \mathcal{G}_r^{\text{deceive}}(\lambda) + (1 - \lambda)(1 - \sigma(\lambda)) \cdot \mathcal{G}_r^{\text{truth}}(\lambda))$, i.e., truth survives and deception survives when unchecked. At any mixing belief λ^* ,

$$C = R \cdot (1 - \lambda^*)\sigma(\lambda^*) + \delta (1 - \lambda^*)\sigma(\lambda^*) \cdot \mathcal{G}_r^{\text{deceive}}(\lambda^*). \quad (\text{OA3.4})$$

Intuition: checking sacrifices future continuation precisely on the deception branch (which it terminates), so the benefit of checking is the contemporaneous compensation R plus the

discounted continuation it forgoes under trust on that branch; at indifference this equals the cost C .

By the same single-crossing logic as in the main text (see Theorem 5.2), best replies are monotone in the public honesty belief: the receiver's optimal inspection is weakly decreasing in λ , so a cutoff λ^* exists; the sender's deception is decreasing in the effective discipline term $(1 - \bar{p}^{\text{check}}(\lambda))$ and in the informativeness of public news (which shapes $\mathcal{G}_s^{\text{truth}} - \mathcal{G}_s^{\text{deceive}}$).

Theorem OA3.1. *Fix $\varepsilon \in (0, 1)$ and $\kappa > 1$. Under bounded payoffs and $\delta \in (0, 1)$, there exists a stationary PBE in the public state λ with a receiver cutoff policy and a (possibly mixed) sender policy. At any mixing belief λ^* , the indifference conditions (OA3.3)–(OA3.4) hold with updates given by (OA3.1)–(OA3.2).*

Proof. Fix $\varepsilon \in (0, 1)$ and $\kappa > 1$. Let $\Lambda := [0, 1]$ denote the compact public belief space. At any $\lambda \in \Lambda$, the sender's and receiver's action sets are finite: $A_s = \{\text{truth}, \text{deceive}\}$ and $A_r = \{\text{trust}, \text{check}\}$. Mixed actions at λ are identified with probabilities in $[0, 1]$, i.e., the sender's deception probability $\sigma(\lambda) \in [0, 1]$ and the receiver's inspection probability $\rho(\lambda) \in [0, 1]$.

Step 1. Let

$$\Sigma := \prod_{\lambda \in \Lambda} [0, 1], \quad P := \prod_{\lambda \in \Lambda} [0, 1], \quad \mathcal{K} := \Sigma \times P.$$

Endow Σ and P with the product topology and $\mathcal{K}$ with the product of these. By Tychonoff, $\mathcal{K}$ is compact; it is convex as a product of convex sets. An element $(\sigma, \rho) \in \mathcal{K}$ is a (possibly nonmeasurable) strategy profile $\lambda \mapsto (\sigma(\lambda), \rho(\lambda))$.

Fix $(\sigma, \rho) \in \mathcal{K}$. Define the public-state value operator $\mathcal{T}_{\varepsilon, \sigma, \rho}$ on the Banach space $\mathcal{V} := \mathcal{B}(\Lambda) \times \mathcal{B}(\Lambda)$ of bounded real-valued functions on Λ (with the sup norm), via the right-hand sides of the Bellman equations in (OA3.3)–(OA3.4): for each $\lambda \in \Lambda$,

$$\mathcal{T}_{\varepsilon, \sigma, \rho} \begin{bmatrix} V_s \\ V_r \end{bmatrix} (\lambda) = \begin{bmatrix} \sigma(\lambda) \left\{ B + \delta [1 - \bar{p}^{\text{check}}(\lambda)] \mathcal{G}_s^{\text{deceive}}(\lambda; \varepsilon, \sigma, V_s) \right\} \\ \quad + (1 - \sigma(\lambda)) \left\{ \delta \mathcal{G}_s^{\text{truth}}(\lambda; \varepsilon, \sigma, V_s) \right\} \\ \\ \rho(\lambda) \left\{ B[\lambda + (1 - \lambda)(1 - \sigma(\lambda))] + \right. \\ \quad R(1 - \lambda)\sigma(\lambda) - C + \\ \delta [\lambda + (1 - \lambda)(1 - \sigma(\lambda))] \mathcal{G}_r^{\text{truth}}(\lambda; \varepsilon, \sigma, V_r) \left. \right\} \\ \quad + (1 - \rho(\lambda)) \left\{ B[1 - (1 - \lambda)\sigma(\lambda)] \right. \\ \quad \quad + \delta [\lambda \mathcal{G}_r^{\text{truth}}(\lambda; \varepsilon, \sigma, V_r) \\ \quad \quad + (1 - \lambda)\sigma(\lambda) \mathcal{G}_r^{\text{deceive}}(\lambda; \varepsilon, \sigma, V_r) \\ \quad \quad \left. + (1 - \lambda)(1 - \sigma(\lambda)) \mathcal{G}_r^{\text{truth}}(\lambda; \varepsilon, \sigma, V_r) \right\} \end{bmatrix}, \quad (\text{OA3.5})$$

where $\bar{p}^{\text{check}}(\lambda) = \mu_0 + (1 - \mu_0)\rho(\lambda)$ and the news-weighted continuation operators $\mathcal{G}_i^{\text{truth}}, \mathcal{G}_i^{\text{deceive}}$ ($i \in \{s, r\}$) are defined in §OA3.1 from the Bayes maps $\lambda^1(\lambda)$ and $\lambda^0(\lambda)$ in (OA3.1)–(OA3.2). For any $(V_s, V_r), (V'_s, V'_r) \in \mathcal{V}$,

$$\left\| \mathcal{T}_{\varepsilon, \sigma, \rho}[V_s, V_r] - \mathcal{T}_{\varepsilon, \sigma, \rho}[V'_s, V'_r] \right\|_{\infty} \leq \delta \left\| [V_s, V_r] - [V'_s, V'_r] \right\|_{\infty},$$

so $\mathcal{T}_{\varepsilon, \sigma, \rho}$ is a contraction with modulus δ . By Banach's fixed-point theorem, for each $(\sigma, \rho) \in \mathcal{K}$ there exists a *unique* bounded value pair

$$(V_s^{\varepsilon, \sigma, \rho}, V_r^{\varepsilon, \sigma, \rho}) \in \mathcal{V}$$

satisfying the public-state Bellman equations.

Step 2. If $(\sigma_n, \rho_n) \rightarrow (\sigma, \rho)$ in the product topology of $\mathcal{K}$, then for each fixed λ the Bayes maps $\lambda^1(\lambda)$, $\lambda^0(\lambda)$ defined by (OA3.1)–(OA3.2) are continuous in $\sigma(\lambda)$, and hence the operators $\mathcal{T}_{\varepsilon, \sigma_n, \rho_n}$ converge pointwise to $\mathcal{T}_{\varepsilon, \sigma, \rho}$ on $\mathcal{V}$. Because the modulus δ is common, the unique fixed points satisfy

$$V_i^{\varepsilon, \sigma_n, \rho_n}(\lambda) \longrightarrow V_i^{\varepsilon, \sigma, \rho}(\lambda) \quad \text{for each } \lambda \in \Lambda, \quad i \in \{s, r\}.$$

Step 3. For any $(\sigma, \rho) \in \mathcal{K}$, define at each $\lambda \in \Lambda$ the sender's and receiver's *pointwise* Bellman objectives under own action $a_s \in A_s$, $a_r \in A_r$ by substituting the unique values $(V_s^{\varepsilon, \sigma, \rho}, V_r^{\varepsilon, \sigma, \rho})$ into the current-payoff plus discounted-continuation formulas (OA3.3)–(OA3.4). Because action sets are finite and payoffs/continuations are continuous in $(\sigma(\lambda), \rho(\lambda))$ and λ , these objectives are continuous in the parameters and *affine* in own mixed action at λ . By Berge's Maximum Theorem, the *pointwise* best-reply correspondences

$$\mathcal{B}_s(\sigma, \rho)(\lambda) \subset [0, 1], \quad \mathcal{B}_r(\sigma, \rho)(\lambda) \subset [0, 1]$$

are nonempty, convex, compact-valued and upper hemicontinuous in (σ, ρ) for each fixed λ .

Define the *global* best-reply correspondence on $\mathcal{K}$ by

$$\mathcal{B}(\sigma, \rho) := \left\{ (\tilde{\sigma}, \tilde{\rho}) \in \mathcal{K} : \tilde{\sigma}(\lambda) \in \mathcal{B}_s(\sigma, \rho)(\lambda), \tilde{\rho}(\lambda) \in \mathcal{B}_r(\sigma, \rho)(\lambda) \quad \forall \lambda \in \Lambda \right\}.$$

Equip $\mathcal{K}$ with the product topology (Tychonoff). Because products of upper hemicontinuous correspondences with nonempty compact convex values are upper hemicontinuous with nonempty compact convex values in the product topology,¹⁰ $\mathcal{B}$ has nonempty, convex, compact values and a closed graph.

Step 4. The Kakutani–Fan–Glicksberg fixed-point theorem applies on the compact convex set $\mathcal{K}$ (a product of intervals; a convex subset of the locally convex product space $\prod_{\lambda \in \Lambda} \mathbb{R}^2$). Therefore there exists $(\sigma^*, \rho^*) \in \mathcal{K}$ with

$$(\sigma^*, \rho^*) \in \mathcal{B}(\sigma^*, \rho^*).$$

By construction, this means that at every $\lambda \in \Lambda$ the strategies $\sigma^*(\lambda)$ and $\rho^*(\lambda)$ are (possibly mixed) pointwise best replies given (σ^*, ρ^*) and their induced values $(V_s^{\varepsilon, \sigma^*, \rho^*}, V_r^{\varepsilon, \sigma^*, \rho^*})$.

Step 5. The public state evolves via the Bayes maps (OA3.1)–(OA3.2), which depend only on current λ and $\sigma^*(\lambda)$; termination after detected deception is absorbing. Hence the value pair $(V_s^{\varepsilon, \sigma^*, \rho^*}, V_r^{\varepsilon, \sigma^*, \rho^*})$ solves the public-state Bellman system under (σ^*, ρ^*) . Because (σ^*, ρ^*) are

¹⁰See, e.g., (Aliprantis and Border, 2006) for closed-graph and u.h.c. stability under products; or argue directly: graphs are closed pointwise and the product of closed sets is closed in the product topology.

pointwise best replies at every λ , the profile is sequentially rational in the public state, and beliefs are Bayes-consistent by construction. Thus (σ^*, ρ^*) constitutes a stationary PBE in the ε -alarm environment.

Step 6. The argument above delivers a fixed point in $\mathcal{K}$; a priori, the selections $\lambda \mapsto \sigma^*(\lambda), \rho^*(\lambda)$ need not be Borel measurable. However, for each (σ, ρ) the pointwise best-reply correspondences $\lambda \mapsto \mathcal{B}_s(\sigma, \rho)(\lambda)$ and $\lambda \mapsto \mathcal{B}_r(\sigma, \rho)(\lambda)$ have measurable graphs (they are closed-valued and depend continuously on λ through continuous payoffs and values). By the (Kuratoski and Ryll-Nardzewski, 1965) measurable selection theorem, there exist Borel-measurable selectors. Therefore we may choose the equilibrium selectors to be Borel-measurable functions of λ .

This completes the proof of existence of a stationary PBE in the public belief state for the ε -alarm model. $\square$

Theorem OA3.2. *Let $(\sigma_\varepsilon, \rho_\varepsilon, \lambda_\varepsilon^*)$ denote stationary PBE objects for each $\varepsilon > 0$. If policies are uniformly Lipschitz on compact subsets of $(0, 1)$ and $\{\lambda_\varepsilon^*\}$ remains in such a subset, then along any sequence $\varepsilon_n \downarrow 0$ there exists a subsequence converging to $(\sigma_0, \rho_0, \lambda_0^*)$ that solves the limit fixed point obtained by replacing the Bayes maps (OA3.1)–(OA3.2) with their $\varepsilon = 0$ limits in (OA3.3)–(OA3.4).*

Proof. Fix $\kappa > 1$ and let $\bar{\varepsilon} \in (0, 1/\kappa)$. Consider any sequence $\{\varepsilon_n\}_{n \geq 1} \subset (0, \bar{\varepsilon}]$ with $\varepsilon_n \downarrow 0$. For each n , let $(\sigma_n, \rho_n, \lambda_n^*)$ be a stationary PBE of the ε_n -alarm environment, where strategies depend on the public belief $\lambda \in (0, 1)$. By assumption, there exist $0 < \underline{\lambda} < \bar{\lambda} < 1$ such that $\lambda_n^* \in [\underline{\lambda}, \bar{\lambda}] =: \Lambda$ for all n , and the strategies are uniformly Lipschitz on Λ .

Step 1. By Arzelà–Ascoli, the sets

$$\mathcal{S} := \{\sigma_n|_\Lambda : n \geq 1\}, \quad \mathcal{R} := \{\rho_n|_\Lambda : n \geq 1\}$$

are relatively compact in $(C(\Lambda), \|\cdot\|_\infty)$ because they are uniformly bounded in $[0, 1]$ and equicontinuous (uniform Lipschitz). Passing to a subsequence (not relabeled), there exist $\sigma_0, \rho_0 \in C(\Lambda)$ with $\sigma_n \rightarrow \sigma_0$ and $\rho_n \rightarrow \rho_0$ uniformly on Λ . Since $\lambda_n^* \in \Lambda$ for all n , compactness yields (again passing to a subsequence if needed) $\lambda_n^* \rightarrow \lambda_0^* \in \Lambda$.

Step 2. For each $\varepsilon \in (0, \bar{\varepsilon}]$, define the Bayes maps at belief λ under ε -alarm (cf. (OA3.1)–(OA3.2)):

$$\begin{aligned} \beta^1(\lambda; \sigma) &:= \frac{\lambda}{\lambda + (1 - \lambda) [\kappa \sigma(\lambda) + (1 - \sigma(\lambda))]}, \\ \beta_\varepsilon^0(\lambda; \sigma) &:= \frac{(1 - \varepsilon)\lambda}{(1 - \varepsilon)\lambda + (1 - \lambda) [1 - \varepsilon(\kappa \sigma(\lambda) + 1 - \sigma(\lambda))]} \end{aligned}$$

For $\lambda \in \Lambda$ and $\sigma \in C(\Lambda)$ with values in $[0, 1]$, the denominators are bounded below by a positive constant independent of $\varepsilon \in (0, \bar{\varepsilon}]$: indeed,

$$\lambda + (1 - \lambda) [\kappa \sigma + (1 - \sigma)] \geq \underline{\lambda} > 0,$$

and, using $\varepsilon \leq \bar{\varepsilon} < 1/\kappa$,

$$(1 - \varepsilon)\lambda + (1 - \lambda) [1 - \varepsilon(\kappa \sigma + 1 - \sigma)] \geq (1 - \bar{\varepsilon})\underline{\lambda} + (1 - \bar{\lambda}) [1 - \bar{\varepsilon} \max\{1, \kappa\}] > 0.$$

Hence the maps $(\varepsilon, \lambda, \sigma) \mapsto \beta^1(\lambda; \sigma)$ and $(\varepsilon, \lambda, \sigma) \mapsto \beta_\varepsilon^0(\lambda; \sigma)$ are jointly continuous on $(0, \bar{\varepsilon}] \times \Lambda \times \mathcal{B}$, where $\mathcal{B}$ is the unit ball of $C(\Lambda)$.

Next define the *news-weighted* continuation operators for $i \in \{s, r\}$:

$$\mathcal{G}_i^{\text{truth}}(\lambda; \varepsilon, \sigma, V_i) := \varepsilon V_i(\beta^1(\lambda; \sigma)) + (1 - \varepsilon) V_i(\beta_\varepsilon^0(\lambda; \sigma)),$$

$$\mathcal{G}_i^{\text{deceive}}(\lambda; \varepsilon, \sigma, V_i) := \kappa \varepsilon V_i(\beta^1(\lambda; \sigma)) + (1 - \kappa \varepsilon) V_i(\beta_\varepsilon^0(\lambda; \sigma)).$$

If $V_{i,n} \rightarrow V_{i,0}$ uniformly on Λ and $\sigma_n \rightarrow \sigma_0$ uniformly, then by the continuity of $\beta^1, \beta_\varepsilon^0$ we have

$$\sup_{\lambda \in \Lambda} |\mathcal{G}_i^{\text{truth}}(\lambda; \varepsilon_n, \sigma_n, V_{i,n}) - \mathcal{G}_i^{\text{truth}}(\lambda; 0, \sigma_0, V_{i,0})| \rightarrow 0,$$

and similarly for $\mathcal{G}_i^{\text{deceive}}$, where we set the *zero-news limits*¹¹

$$\mathcal{G}_i^{\text{truth}}(\lambda; 0, \sigma, V_i) = V_i(\lambda), \quad \mathcal{G}_i^{\text{deceive}}(\lambda; 0, \sigma, V_i) = V_i(\lambda).$$

Step 3. The map $(\varepsilon, \sigma, \rho, V_s, V_r) \mapsto \mathcal{T}_{\varepsilon, \sigma, \rho}[V_s, V_r]$ is jointly continuous on $(0, \bar{\varepsilon}] \times C(\Lambda)^4$ by Step 2 and uniform boundedness of σ, ρ . By a standard parametric contraction argument,¹² if $\varepsilon_n \rightarrow 0$, $\sigma_n \rightarrow \sigma_0$, and $\rho_n \rightarrow \rho_0$ uniformly on Λ , then

$$(V_s^{\varepsilon_n, \sigma_n, \rho_n}, V_r^{\varepsilon_n, \sigma_n, \rho_n}) \longrightarrow (V_s^{0, \sigma_0, \rho_0}, V_r^{0, \sigma_0, \rho_0}) \quad \text{uniformly on } \Lambda.$$

Step 4. For each $\varepsilon \in [0, \bar{\varepsilon}]$ and fixed (σ, ρ) , define at each $\lambda \in \Lambda$ the sender's and receiver's *pointwise* Bellman objectives under own action $a_s \in \{0, 1\}$ ($0 \equiv \text{truth}$, $1 \equiv \text{deceive}$) and $a_r \in \{0, 1\}$ ($0 \equiv \text{trust}$, $1 \equiv \text{check}$) using the corresponding current payoff terms and the continuation given by Step 3. These objectives are continuous in $(\varepsilon, \sigma, \rho, V_s, V_r, \lambda)$ and *affine* in own action. By Berge's Maximum Theorem, the pointwise best-reply correspondences

$$\mathcal{B}_{s, \varepsilon}(\sigma, \rho)(\lambda) \subset [0, 1], \quad \mathcal{B}_{r, \varepsilon}(\sigma, \rho)(\lambda) \subset [0, 1]$$

are nonempty, convex-valued, and upper hemicontinuous in $(\varepsilon, \sigma, \rho)$ (for the product of the sup-norm topologies). Hence the *global* best-reply correspondence

$$\mathcal{B}_\varepsilon(\sigma, \rho) := \left\{ (\tilde{\sigma}, \tilde{\rho}) \in C(\Lambda)^2 : \tilde{\sigma}(\lambda) \in \mathcal{B}_{s, \varepsilon}(\sigma, \rho)(\lambda), \tilde{\rho}(\lambda) \in \mathcal{B}_{r, \varepsilon}(\sigma, \rho)(\lambda) \quad \forall \lambda \in \Lambda \right\}$$

has a closed graph in $C(\Lambda)^2 \times C(\Lambda)^2$.¹³

By stationarity, $(\sigma_n, \rho_n) \in \mathcal{B}_{\varepsilon_n}(\sigma_n, \rho_n)$ for all n . Passing to the limit along the convergent subsequence of Step 1 and using the closed-graph property (and Step 3 for continuity of values), we conclude

$$(\sigma_0, \rho_0) \in \mathcal{B}_0(\sigma_0, \rho_0),$$

¹¹The latter follow from β^1 being independent of ε and $\beta_\varepsilon^0(\lambda; \sigma) \rightarrow \lambda$ as $\varepsilon \downarrow 0$.

¹²E.g., if T_θ are contractions with a common modulus and $T_{\theta_n} \rightarrow T_{\theta_0}$ uniformly on a compact set, then their unique fixed points converge: $V_{\theta_n} \rightarrow V_{\theta_0}$.

¹³Closedness follows from pointwise upper hemicontinuity and Tychonoff's theorem.

i.e., (σ_0, ρ_0) is a stationary best-response pair in the $\varepsilon = 0$ environment (the “no-public-news” limit), with values $(V_s^{0,\sigma_0,\rho_0}, V_r^{0,\sigma_0,\rho_0})$.

Step 5. Let the indifference residuals at ε be

$$\begin{aligned} F_{1,\varepsilon}(\lambda) &:= B - \delta \left[\mathcal{G}_s^{\text{truth}}(\lambda; \varepsilon, \sigma_\varepsilon, V_s^{\varepsilon, \sigma_\varepsilon, \rho_\varepsilon}) - (1 - \bar{p}_\varepsilon^{\text{check}}(\lambda)) \mathcal{G}_s^{\text{deceive}}(\lambda; \varepsilon, \sigma_\varepsilon, V_s^{\varepsilon, \sigma_\varepsilon, \rho_\varepsilon}) \right], \\ F_{2,\varepsilon}(\lambda) &:= C - R(1 - \lambda)\sigma_\varepsilon(\lambda) - \delta(1 - \lambda)\sigma_\varepsilon(\lambda) \mathcal{G}_r^{\text{deceive}}(\lambda; \varepsilon, \sigma_\varepsilon, V_r^{\varepsilon, \sigma_\varepsilon, \rho_\varepsilon}), \end{aligned}$$

where $\bar{p}_\varepsilon^{\text{check}}(\lambda) = \mu_0 + (1 - \mu_0)\rho_\varepsilon(\lambda)$. By construction, at any mixed state λ_ε^* of a stationary equilibrium we have

$$F_{1,\varepsilon}(\lambda_\varepsilon^*) = 0, \quad F_{2,\varepsilon}(\lambda_\varepsilon^*) = 0.$$

By Steps 1–3, $F_{j,\varepsilon_n} \rightarrow F_{j,0}$ *uniformly* on Λ ($j = 1, 2$), where the limit residuals are

$$\begin{aligned} F_{1,0}(\lambda) &= B - \delta \left[V_s^{0,\sigma_0,\rho_0}(\lambda) - (1 - \bar{p}_0^{\text{check}}(\lambda)) V_s^{0,\sigma_0,\rho_0}(\lambda) \right] = B - \delta \bar{p}_0^{\text{check}}(\lambda) V_s^{0,\sigma_0,\rho_0}(\lambda), \\ F_{2,0}(\lambda) &= C - R(1 - \lambda)\sigma_0(\lambda) - \delta(1 - \lambda)\sigma_0(\lambda) V_r^{0,\sigma_0,\rho_0}(\lambda), \end{aligned}$$

with $\bar{p}_0^{\text{check}}(\lambda) = \mu_0 + (1 - \mu_0)\rho_0(\lambda)$. Let $\lambda_n^* \rightarrow \lambda_0^*$ (Step 1). Uniform convergence of F_{j,ε_n} and continuity implies

$$0 = \lim_{n \rightarrow \infty} F_{j,\varepsilon_n}(\lambda_n^*) = F_{j,0}(\lambda_0^*), \quad j = 1, 2.$$

Hence λ_0^* solves the *limit* indifference system obtained by replacing the Bayes maps with their $\varepsilon = 0$ limits in (OA3.3)–(OA3.4), as claimed.

We have extracted a subsequence for which $(\sigma_{\varepsilon_n}, \rho_{\varepsilon_n}, \lambda_{\varepsilon_n}^*) \rightarrow (\sigma_0, \rho_0, \lambda_0^*)$ uniformly on Λ (for policies) and pointwise (for cutoffs), with (σ_0, ρ_0) a stationary best-response pair in the $\varepsilon = 0$ environment and λ_0^* solving the limit indifference equations. $\square$

From (OA3.3), the sender’s margin depends on the gap $\mathcal{G}_s^{\text{truth}} - (1 - \bar{p}^{\text{check}})\mathcal{G}_s^{\text{deceive}}$: larger κ (more informative alarms) raises this gap and thus reduces the deception region; larger ε strengthens news and moves the cutoff similarly; higher C from (OA3.4) raises the honesty cutoff unless offset by higher R ; larger B increases the sender’s temptation and shifts the cutoff in the standard direction.

OA3.2 Self-confirming equilibrium

We now analyze a version of the model in which *both* players’ actions are privately observed and there is *no* public signal (no ε -alarm or leak). A detected deception still terminates the relationship, but detection occurs only when the receiver privately checks and the sender privately deceives. We study *stationary self-confirming equilibria (SCE)* in which players use time-invariant mixing intensities and hold correct beliefs about the distribution of *their own private signals* induced by the strategy profile, without requiring common knowledge of the full play path.¹⁴

¹⁴Our notion matches the standard definition of SCE for repeated interactions with private monitoring: each player’s strategy is a best response to her conjecture about the distribution of signals she observes, and her conjecture is correct *along the realized path*; off-path beliefs need not be correct. See Fudenberg and Levine (1993) for the original definition and Dekel et al. (1999) for a discussion in repeated settings.

Primitives are as in the main text (benefit $B > 0$, cost $C > 0$, discount $\delta \in (0, 1)$, compensation $R \geq 0$ upon detected deception), except there is *no public information*. To allow the receiver to face deception risk even if some senders are intrinsically honest, we keep a (possibly zero) probability $\theta \in [0, 1]$ that the sender is a committed honest type. A strategic sender chooses between **truth** and **deceive**; a strategic receiver chooses between **trust** and **check**. The committed honest sender always plays **truth**; there is no committed vigilant receiver in this subsection (the receiver we study is the long-run strategic player in the relationship).

We focus on *stationary rate profiles*: the strategic sender deceives each period with probability $\sigma \in [0, 1]$, and the receiver checks with probability $\rho \in [0, 1]$. Let

$$\begin{aligned} s &:= (1 - \theta) \sigma && \text{(deception probability in a period),} \\ h &:= \rho s && \text{(termination hazard per period).} \end{aligned}$$

Fix (σ, ρ) and let V_s and V_r denote the (stationary) continuation values at the start of a period for the strategic sender and the strategic receiver, respectively. Because both actions are private and there is no public signal, the only way the relationship ends is if the receiver checks and the sender deceives in that period, which occurs with probability $h = \rho s$.

For the strategic sender, choosing **deceive** yields a current payoff B and continuation only if not checked (probability $1 - \rho$ when he deceives); choosing **truth** yields no current payoff and always continues. Under stationary mixing, the sender's continuation value therefore satisfies

$$V_s = \sigma (B + \delta(1 - \rho) V_s) + (1 - \sigma)(0 + \delta V_s) = \frac{\sigma B}{1 - \delta + \delta \rho \sigma}. \quad (\text{OA3.6})$$

For the strategic receiver, checking yields B on truthful periods and R on deceptive periods, and kills continuation on deception; trusting yields B on truthful periods and 0 on deceptive periods, and never kills continuation. Hence her continuation value under rate ρ is

$$\begin{aligned} V_r &= \rho (B(1 - s) + R s - C + \delta(1 - s) V_r) + (1 - \rho)(B(1 - s) + \delta V_r) \\ &= \frac{B(1 - s) + \rho(R s - C)}{1 - \delta + \delta \rho s}. \end{aligned} \quad (\text{OA3.7})$$

Let $\hat{r}$ denote the sender's conjectured probability of being checked *conditional on deceiving*, and let $\hat{s}$ denote the receiver's conjectured probability of encountering *deception* in a period. In a stationary SCE, conjectures must match the *on-path* frequencies of the private signals the player actually observes:

$$\hat{r} = \rho, \quad \hat{s} = s = (1 - \theta) \sigma \quad \text{if the receiver checks with positive probability;} \quad (\text{OA3.8})$$

if the receiver never checks ($\rho = 0$), then $\hat{s}$ is unrestricted (the receiver observes no informative signal along the path), consistent with SCE.

Definition OA3.3. A stationary SCE is a tuple $(\sigma, \rho, \hat{r}, \hat{s})$ such that: (i) $\sigma \in \arg \max_{\tilde{\sigma} \in [0, 1]} U_s(\tilde{\sigma}; \hat{r})$, where U_s is the sender's stationary value computed from (OA3.6) with ρ replaced by $\hat{r}$; (ii) $\rho \in \arg \max_{\tilde{\rho} \in [0, 1]} U_r(\tilde{\rho}; \hat{s})$, where U_r is the receiver's stationary value computed from (OA3.7) with s replaced by $\hat{s}$; and (iii) the on-path consistency conditions (OA3.8) hold.

The best-response problems are affine in own control because continuation values are linear-fractional but enter the one-shot Bellman comparisons linearly. This yields sharp characterization.

Lemma OA3.4. *Fix any conjectured check rate $\hat{r} \in [0, 1)$ and $\delta \in (0, 1)$. For all $B > 0$, the sender's unique best response is pure deception $\sigma^* = 1$.*

Proof. Given $\hat{r}$, the sender's marginal gain from deceiving today rather than telling the truth equals

$$\Delta_s = B - \delta \hat{r} V_s(\sigma; \hat{r}), \quad V_s(\sigma; \hat{r}) = \frac{\sigma B}{1 - \delta + \delta \hat{r} \sigma}.$$

Evaluating at any $\sigma \in [0, 1]$ gives

$$\Delta_s = B \frac{1 - \delta}{1 - \delta + \delta \hat{r} \sigma} > 0 \quad \text{for } \delta \in (0, 1).$$

Thus deception strictly dominates truth and the unique best response is $\sigma^* = 1$. $\square$

Lemma OA3.5. *Fix any conjectured deception frequency $\hat{s} \in [0, 1]$ and $\delta \in (0, 1)$. The receiver's value under check rate ρ is*

$$V_r(\rho; \hat{s}) = \frac{B(1 - \hat{s}) + \rho(R\hat{s} - C)}{1 - \delta + \delta \rho \hat{s}}.$$

It is (weakly) increasing in ρ iff $R\hat{s} \geq C$ and (strictly) decreasing in ρ iff $R\hat{s} < C$. Consequently,

$$\rho^* = \begin{cases} 1, & \text{if } R\hat{s} > C, \\ [0, 1], & \text{if } R\hat{s} = C, \\ 0, & \text{if } R\hat{s} < C. \end{cases}$$

Proof. Differentiate $V_r(\rho; \hat{s})$ in ρ :

$$\begin{aligned} \frac{\partial V_r}{\partial \rho} &= \frac{(R\hat{s} - C)(1 - \delta + \delta \rho \hat{s}) - \delta \hat{s}[B(1 - \hat{s}) + \rho(R\hat{s} - C)]}{(1 - \delta + \delta \rho \hat{s})^2} \\ &= \frac{(1 - \delta)(R\hat{s} - C)}{(1 - \delta + \delta \rho \hat{s})^2}. \end{aligned}$$

The denominator is positive and $(1 - \delta) > 0$, so the sign is the sign of $(R\hat{s} - C)$, yielding the stated monotonicity and argmax. $\square$

We can now solve for stationary SCE outcomes by combining best responses with on-path consistency.

Theorem OA3.6. *For any $\delta \in (0, 1)$, $B > 0$, $C > 0$, $R \geq 0$, and $\theta \in [0, 1]$, there exists a stationary self-confirming equilibrium. Moreover, in every stationary SCE the strategic sender plays pure deception $\sigma^* = 1$. Let $\hat{s}^* = (1 - \theta)\sigma^* = 1 - \theta$ be the receiver's on-path deception*

risk and let $\hat{r}^* = \rho^*$ be the sender's on-path check rate. Then the receiver's equilibrium checking rate is

$$\rho^* = \begin{cases} 1, & \text{if } R(1 - \theta) > C, \\ [0, 1], & \text{if } R(1 - \theta) = C, \\ 0, & \text{if } R(1 - \theta) < C. \end{cases}$$

In particular, if $R(1 - \theta) \leq C$, the unique stationary SCE outcome has $(\sigma^*, \rho^*) = (1, 0)$ with zero termination hazard and perpetual deception by the strategic sender.

Proof. By Lemma OA3.4, any SCE must have $\sigma^* = 1$. Consistency then implies the receiver's conjectured deception frequency equals the true on-path frequency $\hat{s}^* = 1 - \theta$. By Lemma OA3.5, the receiver's best response is as stated, and we set $\hat{r}^* = \rho^*$ to satisfy on-path consistency for the sender. Thus $(\sigma^*, \rho^*, \hat{r}^*, \hat{s}^*)$ is a stationary SCE for the three cases, yielding existence. Uniqueness of σ^* and the stated characterization of ρ^* follow from the strict inequalities in Lemmas OA3.4–OA3.5. $\square$

From Theorem OA3.6, σ^* is invariant (equal to 1) and ρ^* is weakly increasing in R , weakly decreasing in C , and weakly decreasing in θ (a higher share of intrinsically honest senders reduces the payoff from checking). If $R(1 - \theta) < C$, the receiver never checks, the termination hazard is $h^* = 0$, and the strategic sender's value equals $V_s^* = B/(1 - \delta)$. If $R(1 - \theta) > C$, the receiver always checks, the relationship terminates in the first period unless the sender is intrinsically honest; ex ante the receiver's value is $B\theta + (R - C)(1 - \theta)$ and the strategic sender's value is B (one-period gain), both independent of δ due to immediate absorption.

The SCE analysis shows that with *fully private monitoring and no public information*, dynamic reputational incentives collapse: the strategic sender strictly prefers deception in every period whenever $\delta < 1$, and the receiver's check rate responds only to the static inequality $R(1 - \theta) \geq C$. This degeneracy justifies the minimal public revelation device in Section OA3: even vanishingly rare public news (the ε -alarm with likelihood ratio $\kappa > 1$) restores a Markov public state, stationary PBE existence, and nontrivial cutoff comparative statics in the main variables (B, C, δ) .

OA3.3 Silent-audit leakage $q \in (0, 1)$

We consider a fully private-monitoring environment with *no exogenous public signal*, but where a truthful check is *publicly disclosed* with probability $q \in (0, 1)$ (a “silent audit” that sometimes leaks).¹⁵ The public observes either a disclosure event $D = 1$ (truthful check disclosed), no disclosure $D = 0$, or termination. We show that the public belief about sender honesty λ is Markov, that stationary PBE exist, and that the model nests the main text as $q \uparrow 1$.

At public belief λ , the strategic sender deceives with probability $\sigma(\lambda) \in [0, 1]$ and the receiver inspects with probability $\rho(\lambda) \in [0, 1]$. Let $\bar{p}^{\text{check}}(\lambda) := \mu_0 + (1 - \mu_0)\rho(\lambda)$ denote

¹⁵Termination after detected deception, (deceive, check), is publicly observed and remains absorbing. Receiver checks and sender actions are otherwise private.

the (public) probability that a strategic sender is checked (the receiver may be a committed vigilant type with probability μ_0).

A public disclosure $D = 1$ occurs if and only if a check occurred and truth was verified and disclosed (probability q). Thus

$$\Pr(D = 1 \mid \text{honest}) = q \bar{p}^{\text{check}}(\lambda), \quad \Pr(D = 1 \mid \text{strategic}) = q \bar{p}^{\text{check}}(\lambda) [1 - \sigma(\lambda)].$$

Bayes' rule gives the posterior after disclosure:

$$\lambda^1(\lambda) = \frac{\lambda}{\lambda + (1 - \lambda) [1 - \sigma(\lambda)]}. \quad (\text{OA3.9})$$

Conditional on survival and *no disclosure* $D = 0$, we must rule out observed termination (which would happen only if the strategic sender is checked while deceiving). The probability of $D = 0$ and survival is:

$$\Pr(D = 0, \text{survive} \mid \text{honest}) = 1 - q \bar{p}^{\text{check}}(\lambda),$$

$$\Pr(D = 0, \text{survive} \mid \text{strategic}) = \sigma(\lambda) [1 - \bar{p}^{\text{check}}(\lambda)] + (1 - \sigma(\lambda)) (1 - q \bar{p}^{\text{check}}(\lambda)).$$

Hence the posterior after no disclosure (and no termination) is

$$\lambda^0(\lambda) = \frac{\lambda (1 - q \bar{p}^{\text{check}}(\lambda))}{\lambda (1 - q \bar{p}^{\text{check}}(\lambda)) + (1 - \lambda) (\sigma(\lambda) [1 - \bar{p}^{\text{check}}(\lambda)] + (1 - \sigma(\lambda)) [1 - q \bar{p}^{\text{check}}(\lambda)])}. \quad (\text{OA3.10})$$

The (public) hazard of termination in a period with belief λ equals

$$h(\lambda) = \bar{p}^{\text{check}}(\lambda) (1 - \lambda) \sigma(\lambda), \quad (\text{OA3.11})$$

the probability that a strategic sender deceives and is checked.

Let $V_s(\lambda)$ and $V_r(\lambda)$ be the continuation values when the relationship is active at public belief λ . Define the *disclosure operator* for $i \in \{s, r\}$:

$$\mathcal{G}_{i,q}^{\text{disc}}(\lambda; V_i) := q V_i(\lambda^1(\lambda)) + (1 - q) V_i(\lambda^0(\lambda)). \quad (\text{OA3.12})$$

If the strategic sender plays **deceive** at λ , he receives B and survives only if not checked (probability $1 - \bar{p}^{\text{check}}(\lambda)$), in which case the public sees no disclosure and continues at $\lambda^0(\lambda)$:

$$V_s^{\text{deceive}}(\lambda) = B + \delta (1 - \bar{p}^{\text{check}}(\lambda)) V_s(\lambda^0(\lambda)).$$

If he plays **truth**, he receives 0 and always survives; the next public belief is $\lambda^1(\lambda)$ with probability $q \bar{p}^{\text{check}}(\lambda)$ and $\lambda^0(\lambda)$ otherwise:

$$V_s^{\text{truth}}(\lambda) = \delta \mathcal{G}_{s,q}^{\text{disc}}(\lambda; V_s).$$

Thus

$$V_s(\lambda) = \max \left\{ \delta \mathcal{G}_{s,q}^{\text{disc}}(\lambda; V_s), \quad B + \delta (1 - \bar{p}^{\text{check}}(\lambda)) V_s(\lambda^0(\lambda)) \right\}. \quad (\text{OA3.13})$$

For the receiver, if she checks, her current payoff is $B \cdot (\lambda + (1 - \lambda)(1 - \sigma(\lambda))) + R \cdot (1 - \lambda)\sigma(\lambda) - C$; continuation arises only on truthful periods and equals $\delta \mathcal{G}_{r,q}^{\text{disc}}(\lambda; V_r)$. If she trusts, her current payoff is $B \cdot (1 - (1 - \lambda)\sigma(\lambda))$ and the next public belief is $\lambda^0(\lambda)$ (no disclosure can occur without a check), so continuation is $\delta V_r(\lambda^0(\lambda))$:

$$V_r(\lambda) = \max \left\{ \overbrace{B[1 - (1 - \lambda)\sigma(\lambda)] + \delta V_r(\lambda^0(\lambda))}^{\text{trust}}, \underbrace{B[\lambda + (1 - \lambda)(1 - \sigma(\lambda))] + R(1 - \lambda)\sigma(\lambda) - C + \delta \mathcal{G}_{r,q}^{\text{disc}}(\lambda; V_r)}_{\text{check}} \right\}. \quad (\text{OA3.14})$$

At any mixing belief λ^* , $V_s^{\text{truth}}(\lambda^*) = V_s^{\text{deceive}}(\lambda^*)$ and the two receiver values coincide, yielding:

$$B = \delta \bar{p}^{\text{check}}(\lambda^*) \left[q V_s(\lambda^1(\lambda^*)) + (1 - q) V_s(\lambda^0(\lambda^*)) \right], \quad (\text{OA3.15})$$

$$C = R(1 - \lambda^*)\sigma(\lambda^*) + \delta \left([1 - (1 - \lambda^*)\sigma(\lambda^*)] \mathcal{G}_{r,q}^{\text{disc}}(\lambda^*; V_r) - V_r(\lambda^0(\lambda^*)) \right). \quad (\text{OA3.16})$$

The B terms cancel in (OA3.16) because the expected stage benefit from truth is the same under trust and check.

Theorem OA3.7. *For any $q \in (0, 1)$, $\delta \in (0, 1)$, $B > 0$, $C > 0$, $R \geq 0$, and $\mu_0 \in [0, 1]$, there exists a stationary Perfect Bayesian equilibrium in public state λ with (possibly mixed) policies (σ, ρ) and bounded continuous values (V_s, V_r) solving (OA3.13)–(OA3.14). At any mixed λ^* , the indifference conditions (OA3.15)–(OA3.16) hold with posteriors given by (OA3.9)–(OA3.10).*

Proof. Fix measurable (σ, ρ) . Because $\lambda \mapsto \lambda^1(\lambda)$ and $\lambda \mapsto \lambda^0(\lambda)$ are continuous (denominators are bounded away from zero on $[0, 1]$) and $q \in (0, 1)$, the operator that maps (V_s, V_r) to the right-hand sides of (OA3.13)–(OA3.14) is a contraction with modulus δ on the sup-normed space of bounded functions: the only dependence on (V_s, V_r) is affine through $\mathcal{G}_{i,q}^{\text{disc}}$ and $V_i(\lambda^0(\lambda))$ (each multiplied by δ). By Banach's fixed-point theorem there is a unique bounded continuous value pair (V_s, V_r) for each (σ, ρ) .

Define pointwise best-reply correspondences at each λ by comparing the two affine payoff-continuation expressions for the sender (truth vs. deceive) and the receiver (trust vs. check). Continuity of (V_s, V_r) in (σ, ρ) (parametric contraction) implies these correspondences are nonempty, convex-valued, and upper hemicontinuous in (σ, ρ) (Berge's maximum theorem). On the compact convex strategy space of measurable policies $\lambda \mapsto (\sigma(\lambda), \rho(\lambda)) \in [0, 1]^2$ (product topology), Kakutani–Fan–Glicksberg delivers a fixed point (σ^*, ρ^*) ; Bayes consistency follows by construction from (OA3.9)–(OA3.10). Measurable selection for the pointwise best replies is ensured by Kuratowski–Ryll–Nardzewski. This yields a stationary PBE. $\square$

Proposition OA3.8. *Suppose V_s, V_r are increasing in λ and $\bar{p}^{\text{check}}(\lambda)$ is weakly decreasing in λ (receiver more lenient at higher honesty). Then the sender's best reply is (weakly) decreasing*

in $\bar{p}^{\text{check}}(\lambda)$ and the receiver's best reply is (weakly) decreasing in λ . Consequently, stationary PBE policies admit cutoff characterizations: there exists λ^* such that the receiver checks if and only if $\lambda \leq \lambda^*$, and for each λ the sender deceives with probability decreasing in $\bar{p}^{\text{check}}(\lambda)$ (in particular, pure deception for sufficiently low $\bar{p}^{\text{check}}$).

Proof. The sender's indifference residual at λ equals

$$\begin{aligned}\Delta_s(\lambda) &:= \delta \mathcal{G}_{s,q}^{\text{disc}}(\lambda; V_s) - \left[B + \delta(1 - \bar{p}^{\text{check}}(\lambda)) V_s(\lambda^0(\lambda)) \right] \\ &= -B + \delta \bar{p}^{\text{check}}(\lambda) \left[q V_s(\lambda^1(\lambda)) + (1 - q) V_s(\lambda^0(\lambda)) \right].\end{aligned}$$

This is increasing in $\bar{p}^{\text{check}}(\lambda)$ since V_s is nonnegative; thus the sender is less inclined to deceive when $\bar{p}^{\text{check}}$ is higher. For the receiver, the difference “check – trust” at λ is

$$\Delta_r(\lambda) = R(1 - \lambda)\sigma(\lambda) - C + \delta \left(\left[1 - (1 - \lambda)\sigma(\lambda) \right] \mathcal{G}_{r,q}^{\text{disc}}(\lambda; V_r) - V_r(\lambda^0(\lambda)) \right),$$

which is increasing in λ if V_r is increasing and $\bar{p}^{\text{check}}$ is weakly decreasing in λ (so $\lambda^0(\lambda)$ and $\lambda^1(\lambda)$ are increasing in λ). Hence the receiver is less inclined to check at higher λ , producing a cutoff. $\square$

Theorem OA3.9. *Let $\{q_n\} \subset (0, 1)$ with $q_n \rightarrow q_0 \in (0, 1]$.*

For each n , let $(\sigma_{q_n}, \rho_{q_n}, V_{s,q_n}, V_{r,q_n})$ be a stationary PBE of the silent-audit model with $q = q_n$. There exists a subsequence along which $(\sigma_{q_n}, \rho_{q_n})$ converges uniformly on compact subsets of $(0, 1)$ to a stationary PBE $(\sigma_{q_0}, \rho_{q_0})$ of the model with $q = q_0$, and $V_{i,q_n} \rightarrow V_{i,q_0}$ uniformly on compact subsets for $i \in \{s, r\}$. In particular, as $q \uparrow 1$ the equilibria converge to a stationary PBE of the main termination model where truthful checks are fully public.

Proof. The Bayes maps $\lambda \mapsto \lambda^1(\lambda)$ and $\lambda \mapsto \lambda^0(\lambda)$ depend continuously on (q, σ, ρ) on any compact subset of $(0, 1)$, with denominators bounded away from zero since $q \in (0, 1]$ and $\bar{p}^{\text{check}}(\lambda) \in [0, 1]$. The policy-evaluation operator induced by (OA3.13)–(OA3.14) is a contraction with common modulus δ , jointly continuous in (q, σ, ρ) ; parametric contraction then implies continuity of fixed points (V_s, V_r) in (q, σ, ρ) (Stokey–Lucas–Prescott; Puterman). Pointwise best-reply correspondences are upper hemicontinuous in (q, σ, ρ) (Berge). A standard diagonal extraction (Arzelà–Ascoli on compact subsets using uniform Lipschitz bounds implied by the contraction) yields a subsequence with $(\sigma_{q_n}, \rho_{q_n})$ converging to $(\sigma_{q_0}, \rho_{q_0})$ which is a fixed point of the limiting best-reply correspondence at q_0 (closed-graph argument; Kakutani–Fan–Glicksberg). For $q_0 = 1$, the disclosure operator (OA3.12) collapses to $V_i(\lambda^1(\lambda))$, giving precisely the main model's public-check recursion. $\square$

From (OA3.15), the sender's deterrence margin scales with $q V_s(\lambda^1) + (1 - q) V_s(\lambda^0)$; thus increasing q (greater transparency of truthful checks) raises the right-hand side and reduces the deception region (lower σ or higher receiver cutoff). From (OA3.16), the receiver is more inclined to check when R is larger, C is smaller, or q is higher (truthful checks produce more public discipline). The hazard (OA3.11) obeys the same comparative statics in $\bar{p}^{\text{check}}$, λ , and σ as in the main text.

As $q \downarrow 0$, truthful checks almost never become public; the public belief moves primarily through λ^0 , and dynamic reputational discipline weakens, approaching the SCE benchmark in Section OA3.2. With $q = 1$, silent audits vanish and we recover the main termination model in which truthful checks are publicly observed, so the equilibrium cutoffs and hazards coincide with those derived in the body of the paper (cf. Sections 5 and 8).

OA3.4 Continuous time with Poisson news

We develop a continuous-time counterpart in which public information arrives via a Poisson “news” process that is more likely under deception. Actions are chosen as *rates* that affect the termination hazard and (for the sender) the public news intensity through mixing. We derive the Bayesian filter for the public honesty belief, write the generator, and formulate the stationary HJB system. Under standard regularity (compact controls, bounded intensities), stationary Markov perfect equilibria (MPE) exist; receiver best-responses are cutoff in the public belief, and the model nests the discrete-time ε -alarm limit as $\Delta t \downarrow 0$.

Time is continuous. The discount rate is $\beta > 0$. Let $\lambda_t \in [0, 1]$ denote the public belief at time t that the sender is the committed honest type. A public counting process N_t records “news” arrivals with *state-dependent* intensity: under **truth** the intensity is $\varepsilon > 0$, under **deceive** it is $\kappa\varepsilon$ with $\kappa > 1$ (more alarms under deception).

The sender is either a committed honest type or a strategic type. The committed honest type always chooses **truth**. A *strategic* sender chooses a deception *rate* $\sigma(\lambda) \in [0, 1]$ as a Markov policy. The receiver chooses a private *inspection rate* $r(\lambda) \in [0, \bar{r}]$ with $\bar{r} < \infty$; inspections are not publicly observed. The relationship terminates upon a detected deception, which occurs when the strategic sender is deceiving at that instant and an inspection arrives; the (public) termination intensity at belief λ is

$$h(\lambda) = r(\lambda)(1 - \lambda)\sigma(\lambda). \quad (\text{OA3.17})$$

When termination occurs, the sender’s continuation drops to 0 and the receiver receives a lump-sum compensation $R \geq 0$ contemporaneously (the continuous-time analogue of the discrete compensation).

Conditional on belief λ and strategic mixing $\sigma(\lambda)$, the *observed* public news intensity equals

$$\bar{\lambda}(\lambda) := \varepsilon \left(\lambda + (1 - \lambda) [\sigma(\lambda)\kappa + (1 - \sigma(\lambda))] \right) = \varepsilon \left(1 + (1 - \lambda)\sigma(\lambda)(\kappa - 1) \right). \quad (\text{OA3.18})$$

Let dN_t be the increment of the news process and set the innovation $d\tilde{N}_t := dN_t - \bar{\lambda}(\lambda_{t-})dt$. The (Kallianpur–Striebel/Wonham) filter for point-process observations gives the $\mathbb{F}^N$ -posterior SDE (see, e.g., [Bain and Crisan \(2009, Sec. 9.4\)](#), [Liptser and Shiryaev \(1977, Ch. VI\)](#), [Brémaud \(1981, Ch. VII\)](#)):

$$d\lambda_t = \lambda_{t-}(1 - \lambda_{t-}) \frac{\varepsilon - \varepsilon[1 + \sigma(\lambda_{t-})(\kappa - 1)]}{\bar{\lambda}(\lambda_{t-})} d\tilde{N}_t = -\lambda_{t-}(1 - \lambda_{t-}) \frac{\varepsilon(\kappa - 1)\sigma(\lambda_{t-})}{\bar{\lambda}(\lambda_{t-})} d\tilde{N}_t. \quad (\text{OA3.19})$$

Between news arrivals ($dN_t = 0$), $d\lambda_t = +\varepsilon(\kappa - 1)\sigma(\lambda_t)\lambda_t(1 - \lambda_t)dt$ (belief drifts upward); at an arrival ($dN_t = 1$), λ jumps down to the Bayes posterior

$$\lambda_t = \frac{\lambda_{t-}}{\lambda_{t-} + (1 - \lambda_{t-})[\sigma(\lambda_{t-})\kappa + (1 - \sigma(\lambda_{t-}))]} = \frac{\lambda_{t-}}{1 + (1 - \lambda_{t-})\sigma(\lambda_{t-})(\kappa - 1)}. \quad (\text{OA3.20})$$

Equations (OA3.19)–(OA3.20) define a controlled pure-jump Markov process on $[0, 1]$.

For $f \in C^1([0, 1])$, the (controlled) generator under policy pair (σ, r) is

$$\mathcal{L}^\sigma f(\lambda) = \underbrace{\varepsilon(\kappa - 1)\sigma(\lambda)\lambda(1 - \lambda)f'(\lambda)}_{\text{drift between alarms}} + \underbrace{\bar{\lambda}(\lambda)[f(\lambda^1(\lambda)) - f(\lambda)]}_{\text{downward jump at alarms}}, \quad (\text{OA3.21})$$

where $\lambda^1(\lambda)$ is the post-alarm posterior in (OA3.20). The termination intensity $h(\lambda)$ in (OA3.17) multiplies $-f(\lambda)$ in the HJB (absorption at zero continuation).

Per unit time, the strategic sender obtains flow $B\sigma(\lambda)$; the receiver obtains $B[1 - (1 - \lambda)\sigma(\lambda)]$ (benefit when service is truthful), and pays $C r(\lambda)$; the receiver also obtains a lump-sum R at termination (with intensity $h(\lambda)$).

OA3.4.1 Stationary HJB and equilibrium

For bounded measurable policies (σ, r) , the value functions $V_s, V_r : [0, 1] \rightarrow \mathbb{R}$ solve the linear stationary HJB system

$$\beta V_s(\lambda) = \sup_{\varsigma \in [0, 1]} \left\{ B\varsigma + \mathcal{L}^\varsigma V_s(\lambda) - r(\lambda)(1 - \lambda)\varsigma V_s(\lambda) \right\}, \quad (\text{OA3.22})$$

$$\begin{aligned} \beta V_r(\lambda) = \sup_{\varrho \in [0, \bar{r}]} \left\{ B[1 - (1 - \lambda)\sigma(\lambda)] - C\varrho + \mathcal{L}^\sigma V_r(\lambda) \right. \\ \left. - \varrho(1 - \lambda)\sigma(\lambda)V_r(\lambda) + R\varrho(1 - \lambda)\sigma(\lambda) \right\}, \end{aligned} \quad (\text{OA3.23})$$

with $\mathcal{L}^\varsigma$ denoting (OA3.21) evaluated at the candidate sender rate ς and with boundary conditions $V_s(0) = V_s(1)$ finite, $V_r(0) = V_r(1)$ finite. A *stationary Markov perfect equilibrium (MPE)* is a pair of measurable controls (σ^*, r^*) such that for every λ , $\sigma^*(\lambda)$ attains the supremum in (OA3.22) and $r^*(\lambda)$ attains the supremum in (OA3.23), for the corresponding value pair (V_s, V_r) solving the HJB system.

Theorem OA3.10. *Assume $\beta > 0$, $\varepsilon > 0$, $\kappa > 1$, $B > 0$, $C > 0$, $R \geq 0$, and compact control sets $\varsigma \in [0, 1]$, $\varrho \in [0, \bar{r}]$ with $\bar{r} < \infty$. Then there exists a stationary Markov perfect equilibrium (σ^*, r^*) with bounded continuous value functions V_s, V_r that solve (OA3.22)–(OA3.23).*

Proof sketch. Fix measurable (σ, r) . The state process λ_t is a controlled pure-jump Markov process on the compact metric space $[0, 1]$ with bounded drift and jump intensity (OA3.21). Standard results for controlled Markov processes with bounded rates imply the β -discounted value problem admits a unique bounded solution to the linear HJB (resolvent) equations for *given* (σ, r) (see Davis (1993, Ch. 1–3), Hernández-Lerma and Lasserre (1996, Ch. 10),

Fleming and Soner (2006, Ch. 3) for pure-jump/PDMP settings). Denote the solution $(V_s^{\sigma,r}, V_r^{\sigma,r})$.

Define pointwise best-response correspondences: at each λ , the sender chooses $\varsigma \in [0, 1]$ to maximize the affine map

$$\varsigma \mapsto B\varsigma + \mathcal{L}^\varsigma V_s^{\sigma,r}(\lambda) - r(\lambda)(1-\lambda)\varsigma V_s^{\sigma,r}(\lambda),$$

and the receiver chooses $\varrho \in [0, \bar{r}]$ to maximize the affine map

$$\varrho \mapsto -C\varrho - \varrho(1-\lambda)\sigma(\lambda)V_r^{\sigma,r}(\lambda) + R\varrho(1-\lambda)\sigma(\lambda).$$

Each is a nonempty compact convex-valued correspondence that is upper hemicontinuous in (σ, r) by Berge's maximum theorem (continuity of $\mathcal{L}^\varsigma V_s^{\sigma,r}$ in (ς, σ, r) follows from bounded intensities and continuity of $V_s^{\sigma,r}$; similarly for the receiver). On the compact convex product space of measurable policies $[0, 1]^{[0,1]} \times [0, \bar{r}]^{[0,1]}$ (product topology), the global best-reply correspondence has a closed graph (product of u.h.c. correspondences with compact values). Kakutani–Fan–Glicksberg then yields a fixed point (σ^*, r^*) ; measurable selection is ensured by Kuratowski–Ryll–Nardzewski since the graphs are measurable and values compact. The associated value pair (V_s, V_r) solves (OA3.22)–(OA3.23). $\square$

Lemma OA3.11. *Suppose V_r is increasing in λ and $\sigma(\lambda)$ is weakly decreasing in λ (sender more honest when public believes he is honest). Then the receiver's best response $r^*(\lambda)$ is weakly decreasing in λ and takes a cutoff form: there exists $\lambda^* \in [0, 1]$ with $r^*(\lambda) > 0$ only if $\lambda \leq \lambda^*$; if interior, $r^*(\lambda)$ is uniquely pinned down by the first-order condition*

$$C = (1-\lambda)\sigma(\lambda)(R - V_r(\lambda)). \quad (\text{OA3.24})$$

Proof. For fixed λ , the receiver's choice enters (OA3.23) only through the affine term $-\varrho(1-\lambda)\sigma(\lambda)V_r(\lambda) - C\varrho + R\varrho(1-\lambda)\sigma(\lambda)$. Hence the maximizer is $r^*(\lambda) = \bar{r}$ if the coefficient of ϱ is positive, $r^*(\lambda) = 0$ if negative, and any value if zero. The coefficient equals $(1-\lambda)\sigma(\lambda)(R - V_r(\lambda)) - C$, which is weakly decreasing in λ if V_r is increasing and σ is weakly decreasing in λ ; thus a cutoff exists. If an interior solution is chosen (e.g., by imposing a convex penalty or smoothing), it must satisfy (OA3.24). $\square$

Lemma OA3.12. *Fix $r(\cdot)$. At any λ , the sender's HJB objective in (OA3.22) is strictly decreasing in the effective discipline index*

$$\mathfrak{D}(\lambda; V_s) := r(\lambda)(1-\lambda)V_s(\lambda) - \varepsilon(\kappa-1)\lambda(1-\lambda)\left(\underbrace{V_s(\lambda) - V_s(\lambda^1(\lambda))}_{\text{news gap}}\right).$$

Consequently, if $\mathfrak{D}$ is weakly increasing in λ , then the sender's deception rate $\sigma^(\lambda)$ is weakly decreasing in λ .*

Proof sketch. Differentiate the sender's objective in (OA3.22) with respect to ς using the envelope for the generator: the marginal effect consists of the direct gain B , the marginal

increase in drift $\varepsilon(\kappa - 1)\lambda(1 - \lambda)V'_s(\lambda)$, the marginal increase in news intensity times the jump loss $(V_s(\lambda^1) - V_s(\lambda))$, and the marginal increase in termination hazard $-r(1 - \lambda)V_s(\lambda)$. Grouping terms and using the chain rule $V'_s(\lambda) d\lambda^1/d\varsigma = V_s(\lambda^1) - V_s(\lambda)$ (a standard identity for log-likelihood filters; see [Bain and Crisan \(2009, Prop. 9.4.1\)](#)) gives the stated index $\mathfrak{D}$. Monotonicity then follows. $\square$

Let $\Delta t > 0$ and embed the discrete-time ε -alarm model with period length Δt and discount factor $\delta = e^{-\beta\Delta t}$. Set period alarm probabilities $\varepsilon_\Delta = \varepsilon \Delta t + o(\Delta t)$ and $\kappa\varepsilon_\Delta = \kappa\varepsilon \Delta t + o(\Delta t)$, and let the receiver's check probability be $r(\lambda) \Delta t + o(\Delta t)$ while the sender's deception probability is $\sigma(\lambda) \Delta t + o(\Delta t)$ when using rate controls. Then the discrete Bayes maps and Bellman equations converge to the continuous filter (OA3.19)–(OA3.20), generator (OA3.21), and HJB system (OA3.22)–(OA3.23) as $\Delta t \downarrow 0$ (standard weak-convergence of controlled pure-jump processes; cf. [Ethier and Kurtz \(1986, Ch. 11\)](#)). Higher κ or ε steepens the news gap $V_i(\lambda) - V_i(\lambda^1)$ and the drift in (OA3.21), improving discipline and reducing the deception region; larger B increases the sender's incentives to deceive; higher C depresses inspections via (OA3.24); larger R increases inspections. The hazard $h(\lambda) = r(\lambda)(1 - \lambda)\sigma(\lambda)$ inherits the cutoff tapering in λ from Lemma OA3.11.

The existence proof extends to finite punishment phases by augmenting the state with a finite Markov punishment flag, preserving compactness and bounded rates. Under additional smoothness (e.g., C^1 values), one may express the sender's margin at mixing λ^* in a closed form equating B to a weighted combination of the termination loss $r(1 - \lambda^*)V_s(\lambda^*)$ and the news gap scaled by $\varepsilon(\kappa - 1)$.

References

- Abreu, D., Pearce, D., and Stacchetti, E. (1990). Toward a theory of discounted repeated games with imperfect monitoring. *Econometrica*, 58(5):1041–1063.
- Aliprantis, C. D. and Border, K. C. (2006). *Infinite Dimensional Analysis: A Hitchhiker's Guide*. Springer, Berlin, third edition.
- Bain, A. and Crisan, D. (2009). *Fundamentals of Stochastic Filtering*. Springer, New York.
- Brémaud, P. (1981). *Point Processes and Queues: Martingale Dynamics*. Springer, New York.
- Davis, M. H. A. (1993). *Markov Models and Optimization*. Chapman & Hall, London.
- Dekel, E., Fudenberg, D., and Levine, D. K. (1999). Payoff information and self-confirming equilibrium. *Journal of Economic Theory*, 89(2):165–185.
- Ethier, S. N. and Kurtz, T. G. (1986). *Markov Processes: Characterization and Convergence*. Wiley, New York.

- Fleming, W. H. and Soner, H. M. (2006). *Controlled Markov Processes and Viscosity Solutions*. Springer, New York, 2nd edition.
- Fudenberg, D., Gao, Y., and Pei, H. (2022). A reputation for honesty. *Journal of Economic Theory*, 204:105508.
- Fudenberg, D. and Levine, D. K. (1989). Reputation and equilibrium selection in games with a patient player. *Econometrica*, 57(4):759–778.
- Fudenberg, D. and Levine, D. K. (1993). Self-confirming equilibrium. *Econometrica*, 61(3):523–545.
- Green, E. J. and Porter, R. H. (1984). Noncooperative collusion under imperfect price information. *Econometrica*, 52(1):87–100.
- Hernández-Lerma, O. and Lasserre, J.-B. (1996). *Discrete-Time Markov Control Processes: Basic Optimality Criteria*. Springer, New York.
- Kofman, F. and Lawarrée, J. (1993). Collusion in hierarchical agency. *Econometrica*, 61(3):629–656.
- Kreps, D. M. and Wilson, R. (1982). Reputation and imperfect information. *Journal of Economic Theory*, 27(2):253–279.
- Kuratowski, K. and Ryll-Nardzewski, C. (1965). A general theorem on selectors. *Bulletin of the Polish Academy of Sciences. Mathematics*, 13:397–403.
- Liptser, R. S. and Shiryaev, A. N. (1977). *Statistics of Random Processes I: General Theory*. Springer, Berlin.
- Lukyanov, G. (2023). Collateral and reputation in a model of strategic defaults. *Journal of Economic Dynamics and Control*, 156:104755.
- Lukyanov, G. and Safaryan, S. (2025). Public persuasion with endogenous fact-checking. arXiv:2508.19682. Version dated Aug. 27, 2025.
- Mailath, G. J. and Morris, S. (2002). Repeated games with almost-public monitoring. *Journal of Economic Theory*, 102(1):189–228.
- Marinovic, I., Skrzypacz, A., and Varas, F. (2018). Dynamic certification and reputation for quality. *American Economic Journal: Microeconomics*, 10(2):58–82.
- Marinovic, I. and Szydlowski, M. (2023). Monitor reputation and transparency. *American Economic Journal: Microeconomics*, 15(4):1–67.
- Milgrom, P. and Roberts, J. (1982). Limit pricing and entry under incomplete information: An equilibrium analysis. *Econometrica*, 50(2):443–459.

- Mookherjee, D. and Png, I. (1989). Optimal auditing, insurance, and redistribution. *The Quarterly Journal of Economics*, 104(2):399–415.
- Mookherjee, D. and Png, I. P. L. (1992). Monitoring vis-à-vis investigation in enforcement of law. *American Economic Review*, 82(3):556–565.
- Mookherjee, D. and Png, I. P. L. (1994). Marginal deterrence in enforcement of law. *Journal of Political Economy*, 102(5):1039–1066.
- Townsend, R. M. (1979). Optimal contracts and competitive markets with costly state verification. *Journal of Economic Theory*, 21(2):265–293.