

Screening for Sovereignty: Optimal Procurement when Compliance is Cheap to Fake and Costly to Verify

Stéfane Fermigier · Working paper v0.14 · 8 July 2026

Abstract

Public buyers increasingly attach sovereignty requirements to digital-infrastructure procurement: the supplier must be beyond the reach of foreign compulsion. We model such procurement as mechanism design with a hidden compliance action, a cheap cosmetic-compliance technology, and costly imperfect per-attribute audits backed by a bounded one-strike sanction. Because a detected fake forfeits the same stake however many attributes were faked, violations are complements: per-attribute audit thresholds do not implement compliance, audit costs are superadditive across requirements, and the cost-minimizing audit is a cascade in which each attribute is audited against the stake net of its predecessors' temptations. The price both disciplines the incumbent, as a forfeitable bond, and subsidizes mimicry, by enlarging the prize a mimic keeps when verification fails. Requirements targeting attributes on which supplier types' costs do not differ cannot screen at any price or audit intensity; verification of effective control screens once its accuracy clears a characterized threshold; and when the operational audit merely re-reads the evidence the control audit reads, a single audit is optimal. We characterize when sovereignty mandates are not worth their verification cost, and read the European cloud-procurement record through the results.

Keywords: procurement, costly state verification, multidimensional screening, auditing, digital sovereignty, compliance.

JEL: D82, D86, H57, L51.

Note (preparation and verification). Working paper, not yet peer-reviewed; comments welcome (sf@fermigier.com). The manuscript was prepared with AI assistance under a documented protocol: every closed-form result is verified against independent brute-force computation, statements are cross-read against their proofs at each revision, and the novelty claims were checked adversarially against the literatures engaged in Section 2. The verification suite and the full revision history are available from the author.

1. Introduction

Governments increasingly buy “sovereign” digital infrastructure: cloud, data, and AI services procured under rules meant to guarantee that no foreign power can compel access to the data or withdraw the service. Yet the offerings that satisfy these rules are frequently controlled, in the end, from outside the jurisdiction: an EU-incorporated subsidiary of a non-EU parent, a local data-centre region, a joint venture licensing a foreign partner's technology. The requirement is met on paper while the dependency it was meant to remove is untouched. We call this *sovereignty washing*. It is the motivating instance of a general procurement problem: *how to design requirements that screen for a property the buyer values when that property is multidimensional, hidden, cheap to satisfy nominally, and costly to verify effectively.*

The primitive that organizes the paper is a two-way distinction among the attributes a requirement can target. For each attribute there is an *effective* level (does the supplier actually possess it?) and a *nominal* level (does it appear to?), and the two can be decoupled at a *cosmetic cost*: for jurisdictional control, this is the price of standing up a local subsidiary that does not sever the parent's reach. Attributes then differ in two ways: the *fake-gain* (what a supplier saves by producing the appearance rather than the substance) and the *verifiability* of the gap between the two (the cost and accuracy of

an audit that detects a fake). Legal control is tempting to fake but, being anchored in public records (ultimate ownership, governing law, reach-through exposure), its effective version is cheap and accurate to check. Operational practice is expensive and noisy to check. Data location is easy to check *as a fact* but checking it reveals nothing about the gap that matters, because locating data in-region does not remove it from the reach of whoever controls the operator: for the objective the buyer actually cares about, its verifiability is nil. The integrity of *licensed closed technology* is the still harder case: verifying the absence of compelled covert access in a codebase of hundreds of millions of lines, rewritten with every update, is beyond any buyer's audit capacity at any budget. This attribute's verifiability is nil *by construction*, unless the verification technology itself is changed (source access, reproducible builds) or the attribute is secured structurally (a technology supplier beyond the compelling jurisdiction's reach).

We embed this structure in the simplest mechanism-design model that can carry it. A buyer commits to a requirement set, a price, per-attribute stochastic audits, and a bounded penalty; the supplier privately chooses, attribute by attribute, whether to comply genuinely or cosmetically; a detected fake voids the award and triggers the penalty. Part I of the analysis (Section 4) takes the contracting relationship as given and asks what it costs to make a requirement *credible*; this is a pure hidden-action problem. Part II (Section 5) adds adverse selection: suppliers differ in type (genuinely sovereign or not), the buyer wants the award to go to the right type, or the delivered attributes to come from it, and the question becomes what a requirement can *reveal*.

Four sets of results.

First, fakes are complements, and this changes the economics of multidimensional compliance (Lemma 1, Propositions 1–2). Because a detected fake forfeits the same bounded stake (the price, the operating margin, the legal penalty) no matter how many attributes were faked, the marginal deterrence of an additional audit falls with each attribute already being faked: catching a supplier twice is no worse for it than catching it once. Deterring every attribute *in isolation* therefore does not deter faking them *together*, and the single-attribute audit thresholds familiar from costly-state-verification models are insufficient in the multidimensional problem. The audit intensities that do suffice are superadditive across required attributes: each additional tempting requirement inflates the audit cost of securing every other. Requirements are complements in verification cost. A first design lesson follows: *lean mandates are cheaper than their parts*, since every requirement added to a tender degrades the credibility of the rest unless the audit budget rises more than proportionally.

Second, the cost-minimizing audit is a cascade (Proposition 2, Theorem 1). The overlap in detection events makes spreading audit effort wasteful; the optimum orders the attributes and audits each against the *residual stake*, the stake net of the cumulative temptations of its predecessors, so that only the first attribute sits at its familiar solo-deterrence floor, and every later one is audited strictly harder than in-isolation analysis prescribes. In the two-attribute case this is a corner solution: one attribute at its floor, the entire joint burden loaded on the one with the lowest audit cost per unit of detection, adjusted for temptation. Under the empirically motivated parameterization, in which control is anchored in public records (accurate and cheap to check) while operations are noisy and dear to audit, the loaded attribute is the structural one. *Structural-first verification emerges from the optimization; nothing in the setup assumes it.*

Third, the price is a deterrence instrument, and it cuts in opposite directions on the two margins (Propositions 3 and 5). In the hidden-action problem, the price is a bond: everything the supplier forfeits when caught deters faking, so paying above the participation floor buys deterrence and can be strictly optimal when verification is expensive: lowest-price award rules are anti-compliance by construction. But in the screening problem the same margin flips sign: a higher price is a bigger prize for a *mimic*, and if the audit is imperfect the mimic keeps the prize with positive probability, so

price premia *tighten* the accuracy threshold that verification must clear. Premia are safe on well-verified attributes and subsidize washing on poorly verified ones.

Fourth, and centrally: a requirement screens only where the types' costs diverge and verification reaches (Theorems 2–3). If sovereign and non-sovereign suppliers can both perform an attribute at essentially the same cost (data location, operational practice: the things global suppliers are good at everywhere), then an audit of that attribute, however accurate, detects nothing: the non-sovereign supplier complies with it *genuinely* and fakes only the attribute where the types differ, namely control. If the control attribute itself is unverified, mimicry succeeds at any price, any penalty, and any behavioural-audit intensity; the failure is structural, and no budget repairs it (Theorem 2). Verification of effective control restores screening if and only if its accuracy clears a threshold we characterize in closed form; the threshold is increasing in the mimic's stake in winning and decreasing in the cost of cosmetic compliance, so that corporate-transparency rules (which make shells expensive) and audit accuracy are substitutes (Theorem 3). The requirement-design corollary follows: *require and verify the attribute where the type-cost gap and verifiability are jointly largest, effective control, and require little else; attributes without a type-cost gap contribute nothing to screening and, by complementarity, degrade the credibility of everything else in the tender.*

Because verification is costly, a sovereignty mandate can reduce welfare even when sovereignty is valuable. We characterize the not-worth-it region (Proposition 4) and note that correctly accounting for deviation complementarity *enlarges* it relative to a per-attribute analysis: multidimensional deterrence is more expensive than the sum of its parts, so the corrected model argues *against* mandates on a strictly larger parameter set than a naive version would. The model was built to be able to reach that conclusion.

What we do not claim. Three building blocks are well understood and we use rather than reinvent them. Optimal stochastic auditing with penalties and transfers in a single dimension is the costly-state-verification tradition (Townsend 1979; Border and Sobel 1987; Mookherjee and Png 1989; and for imperfect verification with transfers, Ball and Kattwinkel 2025); our single-attribute floors are exactly that logic. The bare tension between cheap falsification and costly verification is costly state falsification (Lacker and Weinberg 1989), instantiated in insurance-fraud auditing (Dionne, Giuliano and Picard 2009) and quality-indicator gaming (Kuhn and Siciliani 2009). That the price can serve as a forfeitable bond is Becker–Stigler (1974) and Klein–Leffler (1981). Our contribution is what happens when these meet *multidimensionality with attribute-heterogeneous fakeability and verifiability, an endogenous choice of which attributes to require and verify, and a buyer whose objective lives in the attribute that is cheapest to fake and hardest to verify*: the complementarity of fakes, the concentration and superadditivity results, the two-sided role of the price, the impossibility of screening on gap-less attributes, and the accuracy threshold for structural verification. The application to the design of digital-sovereignty procurement is, to our knowledge, the first formal treatment of why structural requirements screen and behavioural ones wash.

Roadmap. Section 2 positions the paper in the literature. Section 3 sets up the model. Section 4 analyzes credibility (hidden action): complementarity, implementability, concentration, the bond, and the not-worth-it region. Section 5 analyzes screening (adverse selection): the washing theorem, the accuracy threshold, the premium inversion, requirement design, and the integrity pole where the optimal mechanism audits only control (Theorem 4). Section 6 confronts the results with the public record. Section 7 discusses robustness and limitations; Section 8 resolves the sufficient-statistic question in both directions and states the one conjecture that remains open; Section 9 concludes. Proofs are in the Appendix.

2. Related literature

Costly state verification. Townsend (1979) and Gale and Hellwig (1985) establish the costly-verification contract; Border and Sobel (1987) and Mookherjee and Png (1989) show that with transfers and penalties optimal auditing is stochastic, targeted at the most tempting misstatements, with penalties at the feasible bound. Ball and Kattwinkel (2025) treat imperfect verification with transfers. Our single-attribute deterrence floors reduce to this tradition, and we say so; the paper departs in the multidimensional structure: joint deviations against a common bounded stake, and the buyer's choice of *which* attributes to require.

Multitask incentives under bounded punishment, and multidimensional enforcement. The supermodularity behind Lemma 1 has a contract-design ancestor: in Bond and Gomes (2009), a multitask agent under limited liability is tempted most by shirking on *all* tasks at once, so that local incentive constraints do not suffice, because bounded pay-for-performance, like our bounded sanction, can only be lost once. Lemma 1 is the audit-design counterpart of that observation, and what we build on it (per-attribute detection probabilities and accuracies, deterrence floors, the residual-stake cascade of Theorem 1, the cost-per-detection allocation, and the superadditivity of audit costs in Corollary 1) has no analogue there. The enforcement literature offers concentration results with different drivers: crackdowns concentrate policing across *agents* (Eeckhout, Persico and Todd 2010), announced monitoring concentrates on few dimensions because *sampling* is coarse (Lazear 2006), and inspection games allocate a detection budget adversarially (Avenhaus, von Stengel and Zamir 2002); marginal deterrence (Shavell 1992; Mookherjee and Png 1994) concerns substitution among offenses of different severity under capped penalties, while our problem is a simultaneous compliance portfolio against a one-strike sanction. Two audit-theoretic neighbours share the multi-front structure with different engines: Rhoades (1999) finds, in multi-component tax reporting, that conditional *sequential* audits create informational spillovers across components and that taxpayers with multiple evasion opportunities may evade more, an echo of our complementarity that runs through information rather than sanction dilution; Patterson and Noel (2003) concentrate the audit on whichever of two frauds the auditee is most motivated to commit, with type-specific penalties doing the ranking. Theirs is concentration by motivation, where ours is forced by detection-overlap against a single bounded sanction, which makes fakes complements rather than alternatives. We found no formal counterpart of the mandate-overload result.

Costly multidimensional screening. Yang (2025) is the nearest neighbour and reaches the opposite headline: with a productive dimension and costly screening instruments, it is optimal (under a positive-correlation condition) to use *no* costly screening and price the productive dimension alone. The reconciliation is structural. Yang's costly instruments are agent-side money-burning that the principal observes freely, and the principal's value lives in the productive, cheaply contractible dimension, so costly instruments are waste to be discarded. Here the principal's value *is* the hidden, hard-to-verify attribute (effective control); there is no productive dimension whose pricing can substitute for verification, and the costly activity is the *principal's* audit, whereas Yang's is the agent's signal. Screening in the costly dimension is therefore the only channel through which the objective can be reached at all.

Costly state falsification and gaming. Lacker and Weinberg (1989), Maggi and Rodríguez-Clare (1995), and Crocker and Morgan (1998) develop falsification at a cost; Dionne, Giuliano and Picard (2009) join falsification to costly audits in insurance fraud; Kuhn and Siciliani (2009) pair adverse selection with manipulation of a quality indicator. These supply our nominal/effective primitive in one dimension. The recent test-design literature (Perez-Richet and Skreta 2022, 2024; Qiu and Shan 2025) studies manipulation-robust test and score design, without transfers, and (in Qiu-Shan) always testing every dimension. Our departures are transfers plus audit-with-penalty plus endogenous

attribute selection. Huang (2026) is the complement: a screening dimension that is verifiable *and non-falsifiable*, where no audit is needed; our problem is the mirror regime where the nominal form is cheaply falsifiable and the audit must target the gap.

Allocation with verification. Ben-Porath, Dekel and Lipman (2014) and Mylovanov and Zapechelnyuk (2017) characterize allocation with costly or ex-post verification and limited penalties, without transfers and in one dimension. We borrow the bounded-penalty discipline; the questions differ as described.

Bonds, prices, and selection. That forfeitable rents deter malfeasance is Becker (1968), Becker and Stigler (1974), Klein and Leffler (1981), Shapiro and Stiglitz (1984); Polinsky and Shavell (2000) survey maximal-sanction results; and in procurement, abnormally low tenders select undercapitalized or reckless suppliers, curable by surety bonds (Calveras, Ganuza and Hauk 2004); this is the discipline face of the price. The selection face has its own ancestors: with a noisy pass/fail test, the stringency needed to keep unqualified applicants from gambling on passing rises with the prize (Guasch and Weiss 1981), which is the screening half of our Proposition 5 in static wage-setting form, and premium prices attract counterfeit mimics under imperfect detection (Grossman and Shapiro 1988). What appears new in Proposition 5 is the *two faces in one mechanism with opposite signs*: the same euro of price relaxes the moral-hazard constraint (Proposition 3) while tightening the required verification accuracy $\alpha_J^*(p)$, yielding the premia-versus-reservation-price design dichotomy. The nearest opposed-sign precedent is Lin (2004), where larger worker-posted bonds relax the no-shirking constraint but degrade the hired pool through a market-level quality externality, an equilibrium composition effect with no verification margin and no accuracy threshold; in Stiglitz and Weiss (1981) both margins worsen with the price, in Bester (1985) collateral helps both, and in Manelli and Vincent (1995) price competition selects bad quality with no verification technology at all. The sign pattern here (help one margin, hurt the other, through the audit-accuracy channel) matches none of them; Calzolari and Spagnolo (2009) trade screening efficiency against relational moral-hazard discipline by restricting the number of bidders and contract duration, with quality observable-but-nonverifiable rather than audited.

Sufficient statistics and monitoring design. With free signals, use them all (Holmström 1979; Amershi and Hughes 1989); with costly acquisition, one signal can suffice for different reasons (Georgiadis and Szentes 2020). Our concentration result (Proposition 2, generalized as the audit cascade in Theorem 1) is driven by neither: it comes from the overlap of detection events against a bounded stake, and its mathematical engine (the incentive constraints carve a contra-polymatroid, so optimal audits are chain-greedy) appears not to have been noted in the auditing literature. A distinct and stronger question, whether with type-linked attributes the *screening* audit can ignore all but the cheapest-to-verify attribute, is settled in both directions in Section 8: false under marginal (Blackwell) domination alone, true when the operational signal is a physical garbling of the control evidence (Theorem 5). This is consistent with Şabac and Yoo (2018): informativeness ranking alone does not license aggregation; a common physical evidence trail does.

Application. The policy motivation draws on weaponized interdependence (Farrell and Newman 2019) and the European debate on cloud sovereignty. We provide the mechanism-design account of why requirements grounded in incorporation or data location fail to screen; the companion policy brief develops the institutional detail.

3. The model

3.1 Primitives

A *buyer* procures one contract. A *supplier* (the firm) can perform the underlying service at price p we normalize the firm's service-delivery cost into p (so p is the net price, the quasi-rent on the contract)

and let $\pi \geq 0$ denote the *ancillary rents* the firm earns from holding the contract *beyond* the price (follow-on business, ecosystem and lock-in value, reference effects), all forfeited if the award is voided. There is no asymmetric information about π in Part I; in Part II the types' ancillary rents differ, and this is where the empirically relevant asymmetry lives. The buyer values the service (at v) and, separately, values *effective sovereignty*.

Sovereignty has n attributes, indexed i the leading case is $n = 2$ with $i \in \{J, O\}$: J is structural/jurisdictional control (who ultimately controls the operator, whose law reaches it), O is operational practice. For each attribute the firm chooses an *effective* level $s_i \in \{0, 1\}$ and presents a *nominal* level $\hat{s}_i \in \{0, 1\}$, $\hat{s}_i \geq s_i$. Eligibility requires $\hat{s}_i = 1$ on every attribute in the buyer's *required set* R . Genuine compliance ($s_i = \hat{s}_i = 1$) costs $C_i > 0$ *cosmetic compliance* ($\hat{s}_i = 1, s_i = 0$) costs $\varphi_i \in [0, C_i)$. The *fake-gain* is

$$\Delta_i = C_i - \varphi_i > 0,$$

what the firm saves by producing the appearance without the substance. For J , φ_J is the cost of the shell (incorporating locally, staffing a board), while C_J is the cost of genuinely severing control; Δ_J is large.

The buyer can *audit* attribute i with probability $a_i \in [0, 1]$ at cost k_i per audit. An audit of a faked attribute detects the fake with probability $\alpha_i \in [0, 1]$ an audit of a genuine attribute never (falsely) flags it. We write

$$x_i = a_i \alpha_i \in [0, \alpha_i]$$

for the per-attribute detection probability, and $c_i = k_i / \alpha_i$ for the audit outlay per unit of detection probability. *The accuracy α_i measures the detectability of the nominal/effective gap; the checkability of the raw state is a different object.* This is where “data residency” gets its formal treatment: whether data sits in-region is trivially checkable, but that check carries no information about whether the data is beyond foreign reach. For the sovereignty objective, residency's $\alpha \approx 0$ while its $\Delta > 0$.

Detection of a fake on any required attribute *voids the award*: the firm loses the price and the margin and pays a penalty $L \leq \bar{L}$ (limited liability: the total sanction, forfeiture plus fine, is bounded by the firm's stake and the legal maximum; detection of several fakes triggers the same voiding and capped penalty). Define the *stake*

$$\Omega = p + \pi + L.$$

Timing: (1) the buyer publicly commits to the mechanism $(R, p, \{a_i\}, L)$ (2) the firm accepts or rejects; (3) if it accepts, it privately chooses which required attributes to fake, a set $D \subseteq R$, paying φ_j for $j \in D$ and C_i for $i \in R \setminus D$ (4) audits run independently per attribute; (5) if no fake is detected the price is paid and the service delivered; otherwise the award is void and L is paid.

The buyer's payoff if the contract completes is $v + \omega \cdot \mathbf{1}\{s_i = 1 \ \forall i \in R^*\} - p - \sum_i a_i k_i$, where R^* is the set of attributes on which effective sovereignty depends and ω the value of sovereignty; if sovereignty fails the buyer bears a loss K (the expected cost of compelled access, denial, and lock-in: the weaponized-interdependence exposure). Section 4 studies the cost side (implementing compliance cheaply); Section 4.5 puts the pieces into the buyer's full comparison.

Everyone is risk-neutral. The buyer commits. *Mechanism class.* We study *procurement mechanisms*: the buyer may post menus of contracts, charge application fees, and randomize the award, but *transfers flow only to the awarded supplier* (winner-only transfers). This is the institutionally relevant class, since public buyers cannot lawfully pay losing bidders, and it is also the robust one: a mechanism that paid non-winners their mimicry rents to stand aside would unravel under free entry

of non-sovereign applicants, each collecting the stand-aside payment. Within this class the restriction to the posted mechanisms above is without loss for our questions (Theorem 2(a) establishes the Part II invariance to menus, fees, and randomization; the hidden action enters through the obedience constraints analyzed directly). All parameters are common knowledge except the firm's choices (Part I) and, in Part II, its type.

The unraveling claim closes the Myerson escape route (paying the mimic to stand aside), so it deserves a formal statement; a skeptical reader may also ask whether a budget-balanced fee scheme evades it.

Remark 0 (winner-only transfers are the robust class). *Extend the mechanism space to allow transfers to non-winning applicants, and let each applicant bear an application cost $\kappa \geq 0$. (a) If in equilibrium some participation strategy followed by losing carries a net expected transfer above κ , the mechanism unravels under free entry: an unbounded pool of non-sovereign firms can replicate that strategy's observable actions (reports and fee payments are replicable by construction, and a non-winner delivers nothing and is never audited, so its type is never engaged), and the buyer's outlay on non-winners grows without bound in the number of entrants. In any free-entry equilibrium, non-winning strategies therefore carry no rents. (b) Free-entry-proof schemes with transfers to non-winners (entry fees with rebates to losers, however structured, with net transfers to non-winners nonpositive) screen no better than winner-only mechanisms: both types face the identical fee-and-rebate stream on identical reports, so the shadow-strategy coupling of Theorem 2(a) applies verbatim and the mimicry wedge is unchanged. (Proof in Appendix.)*

3.2 Remarks on the modeling choices

(i) *Binary levels are without loss under detection-in-kind.* Suppose compliance were continuous, with the firm choosing a shortfall on each attribute, savings increasing in the shortfall, and audits detecting *whether* an attribute falls short (probability α_i for any positive shortfall) rather than *how much*.

Lemma 0 (all-or-nothing fakes). *Under detection-in-kind, every best response of the firm sets each attribute's shortfall to 0 or its maximum: conditional on the set of attributes that fall short, the detection distribution is fixed while savings strictly increase in each shortfall. The binary model is therefore without loss of generality.*

The premise is the right one for structural facts: a shell either severs control or it does not; reach-through either exists or it does not. Where detection plausibly scales with the size of the gap (operational attributes), the continuum problem is genuinely different and open (Conjecture 2). (ii) *No false positives* stacks the deck in favour of auditing; even so, verification will turn out to be expensive or useless in exactly the configurations that matter. (iii) *Bounded total sanction* is the realistic case in public procurement: the operative sticks are award voiding, exclusion, and capped damages; it is also what makes the multidimensional problem non-trivial (an unbounded per-count fine would let per-attribute analysis go through). (iv) The cosmetic technology and the audit accuracies are primitives here; the political economy of *who chooses α* (incumbents lobbying the certification scheme that fixes the audit's reach) is a distinct question we develop in a companion paper, and the choice of R by the buyer is analyzed in 5.3.

4. Part I. Credibility: what it costs to make a requirement real

Throughout this section the firm's type is known (the interesting case: an opportunistic firm, tempted by every Δ_i); the problem is pure hidden action. The buyer wants full genuine compliance on R we ask when that is implementable and at what audit cost. Take $R = \{J, O\}$ unless stated; general- n statements are flagged.

The firm's payoff from accepting and faking the set $D \subseteq R$ is

$$u(D) = \Omega S(D) - L - \sum_{i \in R} C_i + \sum_{j \in D} \Delta_j, \quad S(D) = \prod_{j \in D} (1 - x_j),$$

with $S(\emptyset) = 1$, so $u(\emptyset) = p + \pi - \sum_{i \in R} C_i$. Full compliance is incentive-compatible iff $u(\emptyset) \geq u(D)$ for all D :

$$\Omega [1 - S(D)] \geq \Delta(D) \equiv \sum_{j \in D} \Delta_j \quad \text{for every } D \subseteq R.$$

The single-attribute constraints are the familiar CSV floors: $x_j \geq \delta_j \equiv \Delta_j / \Omega$: audit probability equal to temptation over stake, as in Border–Sobel and Mookherjee–Png. The multidimensional content is everything beyond them.

4.1 Fakes are complements

Lemma 1 (complementarity of fakes). *The gain from adding attribute j to a fake set D , namely $\Delta_j - \Omega x_j S(D)$, is increasing in D . Consequently, if every single-attribute constraint holds with equality ($x_j = \delta_j$ for all j), every joint deviation $|D| \geq 2$ is strictly profitable:*

$$\Omega [1 - \prod_{j \in D} (1 - \delta_j)] = \Delta(D) - V(D) < \Delta(D), \quad V(D) \equiv \Omega \left[\sum_{j \in D} \delta_j - 1 + \prod_{j \in D} (1 - \delta_j) \right].$$

The overlap value $V(D)$ is the amount by which the sum of the solo detection probabilities exceeds the probability of their union, priced at the stake; it is strictly positive for every $|D| \geq 2$ with interior floors (for a pair, $V(\{i, j\}) = \Omega \delta_i \delta_j$).

The economics: the sanction is bounded and indivisible, so a firm caught on two attributes loses exactly what a firm caught on one loses. Detection events overlap, so the *marginal* deterrence of each additional audit falls with every attribute already being faked, while the *marginal* gain Δ_j does not. Faking, like sinning against a one-strike rule, exhibits economies of scale. The bounded-*total*-sanction structure is what drives this: if instead each detected fake triggered its own unbounded fine, deterrence would be additive and per-attribute analysis exact; with per-count fines capped in total (the procurement reality: an award can only be voided once, debarment happens once, damages are capped by the firm's stake), the complementarity survives and degrades continuously as per-count sanctioning capacity grows (Proposition 8 in Section 7 displays the one-parameter family and the exact deficit). Two consequences follow immediately. First, the per-attribute CSV thresholds, under which each attribute is deterred “in isolation”, do *not* implement multidimensional compliance; analyses that treat each requirement's enforcement as a separate auditing problem overstate what a given audit budget achieves. (The contract-design ancestor of this observation is Bond and Gomes 2009, where the binding deviation is shirking on all tasks at once; see Section 2.) Second, incentive compatibility is governed by the *family* (IC- D), and the binding constraint at the cost-minimizing audit is the *joint* deviation.

4.2 What can be secured at all

Proposition 1 (implementability). (a) *At given (p, L) , full compliance on R is implementable iff for every $D \subseteq R$:*

$$\Delta(D) \leq \Omega [1 - \prod_{j \in D} (1 - \alpha_j)].$$

(b) *Allowing the price to be chosen freely, full compliance on R is implementable at some finite price iff $\alpha_i > 0$ for every $i \in R$ with $\Delta_i > 0$. (c) If $\alpha_i = 0$ and $\Delta_i > 0$ for some required attribute, then in*

every mechanism the firm fakes i : washing on i is the unique outcome, at any price, audit budget, and penalty.

Part (c) is the formal version of the residency critique: when the effective version of a requirement cannot be verified *at all*, the requirement is fake with certainty, whatever the budget, and requiring it adds temptation (Δ_i) to the tender while adding no detectable surface. Part (a) contains a capacity constraint: at a given stake, the *sum* of fake-gains that can be deterred is bounded even by perfect audits ($\Delta(R) \leq \Omega$ when all $\alpha_j = 1$), because the stake is a common deterrence resource that all requirements draw on. Raising the price relaxes it (that is Proposition 3); accuracy failures cannot be priced away (that is part (c)).

4.3 The cost-minimizing audit: concentration

Assume implementability with headroom: $\bar{\gamma} \equiv 1 - \delta_J - \delta_O \in (0, 1)$. The buyer solves

$$\min_{x_J, x_O} c_J x_J + c_O x_O \quad \text{s.t.} \quad x_i \geq \delta_i, \quad (1 - x_J)(1 - x_O) \leq \bar{\gamma}, \quad x_i \leq \alpha_i.$$

Proposition 2 (concentration). *The joint constraint binds at the optimum ($(1 - x_J)(1 - x_O) = \bar{\gamma}$), and along it the audit cost is strictly concave, so the optimum is a corner: one attribute is held at its solo floor δ_i and the other carries the entire joint-deterrence burden. Writing $\eta_i \equiv c_i(1 - \delta_i)$, the loaded attribute is the one with the smaller η_i (assume interior feasibility, $x_i^{\text{load}} \leq \alpha_i$ otherwise the cap binds and the remainder spills to the other attribute). In the load- J regime the optimal profile and cost are*

$$x_O^* = \delta_O, \quad x_J^* = \frac{\delta_J}{1 - \delta_O}, \quad A = c_J \frac{\Delta_J}{\Omega - \Delta_O} + c_O \frac{\Delta_O}{\Omega}.$$

Three readings. (i) *Why concentrate*: spreading audits across attributes maximizes the overlap of detection events, the states in which two audits both catch a firm that can only be punished once. Concentration minimizes wasted overlap. (ii) *The index*: $\eta_i = (k_i/\alpha_i)(1 - \Delta_i/\Omega)$ orders attributes by audit outlay per unit of detection, discounted by temptation. Cheap-and-accurate verification (low k_i , high α_i) attracts the load; for sovereignty, control facts are public-record (c_J small) while operational audits are dear and noisy (c_O large), so *the structural attribute optimally carries the joint burden; the model derives structural-first instead of assuming it*. (iii) *The cross-shadow*: in (4.1) the loaded attribute's intensity is its own floor *inflated* by $1/(1 - \delta_O) > 1$: the mere presence of another tempting requirement makes the first one costlier to secure.

Corollary 1 (mandate overload). *The minimal audit cost of a credible requirement set is superadditive: for $R = \{J, O\}$, $A > c_J \delta_J + c_O \delta_O$ (the sum of standalone costs at the same stake), and adding a requirement with $\Delta > 0$ strictly raises the audit cost of securing the existing one. Theorem 1 gives the general- n form: every attribute's audit is inflated by the full temptation of its cascade predecessors.*

Requirements are complements in verification cost. The design lesson is anti-checklist: every clause added to a sovereignty tender (certifications, staffing rules, reporting duties) either brings its own audit budget or dilutes the credibility of the clauses that matter.

The two-attribute corner solution generalizes, in a form that per-attribute intuition does not suggest:

Theorem 1 (the audit cascade). *In the n -attribute problem, under interior feasibility (accuracy caps slack), the cost-minimizing audit is a **cascade**: there is an ordering σ of the attributes such that, writing $S_{k-1} = \{\sigma(1), \dots, \sigma(k-1)\}$,*

$$x_{\sigma(k)} = \frac{\Delta_{\sigma(k)}}{\Omega - \Delta(S_{k-1})}, \quad k = 1, \dots, n,$$

each attribute audited at its deterrence floor computed against the residual stake (the stake net of the cumulative fake-gains of its predecessors), and the optimal ordering minimizes the resulting cost $\sum_k c_{\sigma(k)} x_{\sigma(k)}$. Only the first attribute sits at its solo floor; every later attribute is audited strictly harder than in-isolation analysis prescribes. For $n = 2$ this is Proposition 2, and the optimal ordering is given by the η -index; for $n \geq 3$ no prefix-free index exists in general (the adjacent-exchange comparison, $c_a(\Omega - \Delta(S) - \Delta_a) \geq c_b(\Omega - \Delta(S) - \Delta_b)$, depends on the prefix), so ordering the cascade is a genuine sequencing problem.

The proof (Appendix) exposes the structure behind everything in this section: in hazard coordinates the incentive family (IC- D) carves a *contra-polymatroid* (the constraint function is a convex transform of an additive one, hence supermodular) whose extreme points are exactly the chain-greedy vectors, and audit cost is concave, so the optimum is one of them. The economics of the formula: *deterrence capacity is consumed sequentially*. Attribute k faces not the full stake Ω but what its predecessors' temptations have left of it; as a mandate's total temptation $\Delta(R)$ approaches the stake, the marginal requirement's audit cost explodes like $1/(\Omega - \Delta(R \setminus \{n\}))$: this is Corollary 1's overload, in closed form. Figure 1 displays the cascade for $n = 3$.

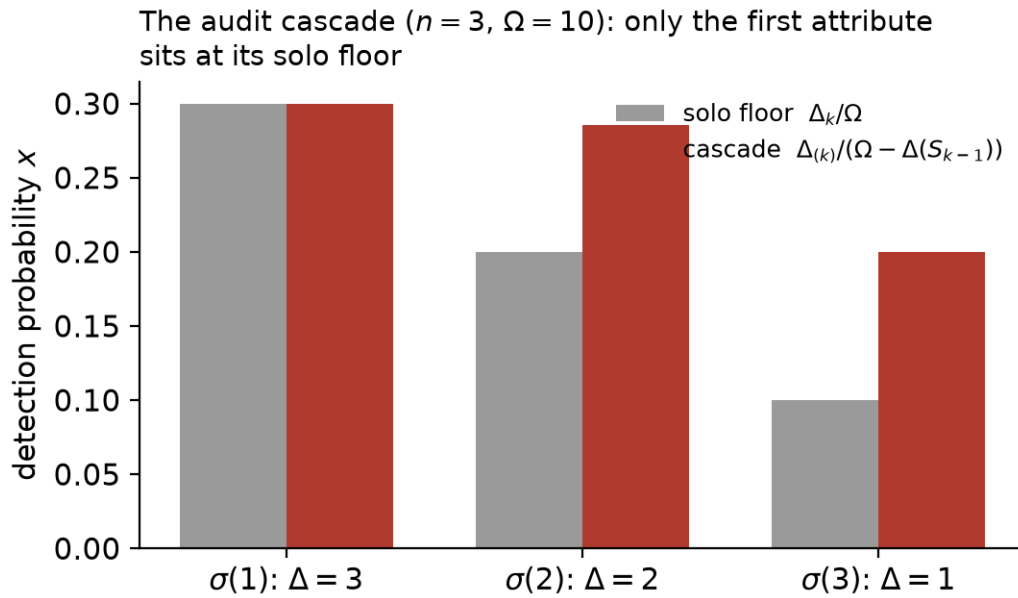


Figure 1. The audit cascade (Theorem 1) for $n = 3, \Omega = 10, \Delta = (3, 2, 1)$. Grey bars are the per-attribute solo floors Δ_k/Ω red bars are the cascade intensities $\Delta_{(k)}/(\Omega - \Delta(S_{k-1}))$. Only the first attribute sits at its solo floor; each later attribute is audited against the residual stake, hence strictly harder.

A numerical illustration. Take a stake of $\Omega = 10$ and requirements each carrying fake-gain $\Delta_i = 1$ and unit cost-per-detection $c_i = 1$. Audited in isolation, each requirement needs $x = 0.10$ the cascade requires $\sum_{k=0}^{n-1} 1/(10 - k)$ in total, against $n/10$ for the naive budget. At $n = 3$ the naive per-attribute budget understates the true cost by 12%; at $n = 6$ by 41%; at $n = 9$ by 114%, so more than half the true bill is missing (Figure 2). A tenth requirement would need certain detection ($x = 1$), and an eleventh is infeasible at any audit intensity: the mandate's total temptation has consumed the stake.

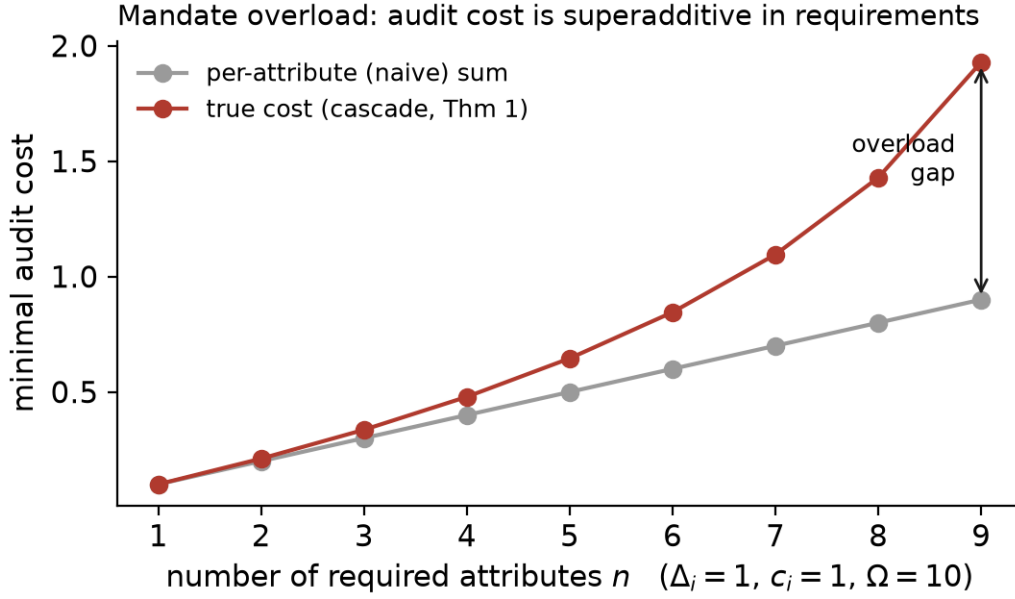


Figure 2. Mandate overload (Corollary 1). Minimal audit cost of a credible mandate as the number of required attributes grows ($\Delta_i = 1$, $c_i = 1$, $\Omega = 10$): the naive per-attribute sum against the true cascade cost. The gap is the joint-deviation burden that per-attribute analysis omits.

4.4 The price as a bond

The stake $\Omega = p + \pi + L$ is an instrument: everything the firm forfeits when caught is deterrence capital. Since A in (4.1) is decreasing and convex in Ω , the buyer's total outlay $T(p) = p + A(p)$ is convex, and:

Proposition 3 (optimal premium). Suppose $c_J \leq c_O$ and $\Delta_J \geq \Delta_O$, so that $\eta_J \leq \eta_O$ at every price and the load- J regime of Proposition 2 holds along the entire price path. Set $L = \bar{L}$ (maximal penalty is costless deterrence). The buyer's optimal price satisfies $p^* > \underline{p}$ (the participation floor) iff $-A'(\underline{p}) > 1$, i.e. iff

$$\frac{c_J \Delta_J}{(\underline{p} + \pi + \bar{L} - \Delta_O)^2} + \frac{c_O \Delta_O}{(\underline{p} + \pi + \bar{L})^2} > 1,$$

in which case p^* solves the first-order condition $-A'(p^*) = 1$: the last euro of price premium saves exactly one euro of audits.

When verification is expensive (large c_i) or the legal penalty weak (small $\bar{L}$), the buyer optimally pays *above* the lowest acceptable price and lets the firm's fear of forfeiting the margin do the auditing's work: the Becker–Stigler bond, inside a tender. The policy translation is direct: *lowest-price award rules are anti-compliance by construction*, because they strip the winner of exactly the rent whose loss would have disciplined it. (Section 5 will qualify this sharply: the bond logic holds for disciplining the firm you have; it inverts for screening the firm you choose.)

4.5 When the mandate is not worth it

Let $V = \omega$ be the value of effective sovereignty, K the exposure borne if it fails, ρ the price premium of the compliant solution over the unconstrained one, and A the audit cost (4.1). Comparing the buyer's three options (mandate-and-verify, buying from a genuinely sovereign entrant at a quality discount, no mandate):

Proposition 4 (the not-worth-it region). *No mandate is optimal iff $V + K < \rho + A$ the entrant dominates the audited incumbent iff the quality gap is smaller than the audit-plus-premium differential. Both regions are non-empty on an open set of parameters, and both are strictly larger than a per-attribute (single-deviation) analysis would compute, because A embeds the joint-deviation burden of Lemma 1.*

Two implications. Verification cost alone can make a sovereignty mandate welfare-reducing even when sovereignty is genuinely valuable, and the correction in this paper (fakes are complements) makes that region strictly *bigger*. And a credible domestic entrant is valuable to the buyer even if never chosen: it caps the price-and-audit bill the incumbent route can command (the outside option economizes on verification).

5. Part II. Screening: what a requirement can reveal

Now let the firm's type be private: $\theta \in \{G, B\}$. G is genuinely sovereign (it possesses the control attribute, or can hold it at low cost); B is not. Costs $C_i(\theta)$, with the same cosmetic technology φ_i for both. Three assumptions carry the analysis, which we take to be *the empirically relevant configuration* rather than conveniences:

- (A1) *Authenticity.* $C_J(G) \leq \varphi_J$: for the genuine type, being what it is costs no more than counterfeiting it. (An EU-controlled operator does not need a shell.)
- (A2) *Behavioural parity.* $C_O(B) = C_O(G) = C_O$: operational practice costs the types the same. (Global suppliers run data centres and staff support desks at least as well as anyone; operations carry no information about control.)
- (A3) *Behavioural discipline is feasible.* $\alpha_O \geq \Delta_O/\Omega_G$ at the relevant prices, with $\Delta_O = C_O - \varphi_O$: the winner's operational moral hazard can be deterred at the Part I floor. (Without it, even the genuine type's contract cannot be made credible on O , a Part I problem logically prior to screening; the assumption is maintained so that the theorems below isolate the screening question.)

The types' ancillary rents are π_G and $\pi_B = \pi_G + e$ with *edge* $e \geq 0$: the non-sovereign incumbent typically has more to gain from holding the contract (ecosystem, lock-in, reference value), and this is what makes the problem hard. For the sharpest statements we take the polar case where B does not genuinely cede control at any admissible price ($C_J(B)$ prohibitive); Remark 4 relaxes this. B 's honest option is to exit (payoff 0). *Washing* means: the award goes to a firm with $\hat{s}_J = 1, s_J = 0$.

The buyer posts $(R = \{J, O\}, p, x, L)$ both types can apply. G 's participation requires $p + \pi_G \geq C_J(G) + C_O$, so the lowest admissible price is $\underline{p} = C_J(G) + C_O - \pi_G$ (assumed positive). If B applies it fakes J (by the polar assumption) and chooses whether to also fake O its deviation payoffs, with $\Omega_B = p + \pi_B + L$:

$$u_B(\{J\}) = \Omega_B(1 - x_J) - L - \varphi_J - C_O, \quad u_B(\{J, O\}) = \Omega_B(1 - x_J)(1 - x_O) - L - \varphi_J - \varphi_O.$$

A *no-washing mechanism* is one G accepts and under which both are ≤ 0 .

5.1 The washing theorem

Theorem 2 (gap-less attributes cannot screen). *Suppose the control attribute is unverified: $\alpha_J = 0$, as when the tender checks location, staffing, and certifications while leaving effective control unexamined. Define the **mimicry wedge** $w \equiv e - (\varphi_J - C_J(G))$. Then:*

(a) *The wedge is mechanism-invariant: in every procurement mechanism (any menu of contracts, application fees, award randomization: any winner-only-transfer scheme), B 's payoff from replicating G 's strategy while faking J exceeds G 's payoff from that same strategy by exactly w , independent of the price, the penalty, and the behavioural-audit intensities.*

(b) If $w > 0$, no procurement mechanism procures genuine sovereignty: every equilibrium either awards, with positive probability, a contract whose control requirement is faked (washing), or never awards a sovereignty-requiring contract at all.

(c) If $w \leq 0$ (and (A3) holds), a no-washing mechanism exists: the floor-price contract with the Part I operational audit floor.

The intuition is the paper's core. With control unverified, B 's cheapest mimicry is to fake J and comply genuinely with everything the buyer *can* check; by behavioural parity, genuine operational compliance costs B exactly what it costs G . The audit therefore never fires: *you cannot catch a mimic on the attributes it isn't faking*. The comparison that remains is pure economics: B 's edge e against the only cost difference between the types that the mechanism engages, $\varphi_J - C_J(G)$, the shell's cost over authenticity. Part (a) is why no cleverness in mechanism design escapes this: whatever bundle of price, audits, fees, and lotteries the buyer offers, B can consume the identical bundle at cost difference $-w$, so the wedge rides through every instrument. (Paying B to stand aside would beat it; Remark 0 is why that escape is excluded, by law and by free entry.) In the application, e is measured in ecosystem and lock-in rents and $\varphi_J - C_J(G)$ in legal fees; $w > 0$ by orders of magnitude. *Data-residency and operational-practice regimes go beyond screening weakly: they select precisely for the suppliers best at operations regardless of control, which is the observed outcome.*

5.2 Structural verification and the accuracy threshold

Theorem 3 (screening on control). Let $\alpha_J > 0$, maintain (A3), and set the price at the participation floor $\underline{p}$ and $L = \bar{L}$. A no-washing mechanism exists iff

$$\alpha_J \geq \alpha_J^* \equiv 1 - \frac{\bar{L} + \varphi_J + C_O}{\underline{p} + \pi_B + \bar{L}} \quad (\text{fake-}J \text{-only constraint}),$$

together with the joint-fake analogue $(1 - \alpha_J)(1 - x_O) \leq (\bar{L} + \varphi_J + \varphi_O) / (\underline{p} + \pi_B + \bar{L})$, which the (A3) discipline floor on x_O relaxes. The threshold α_J^* is:

- increasing in the mimic's stake in winning (π_B , hence in the edge e) and in the price $\underline{p}$
- decreasing in the shell cost φ_J and in the penalty cap $\bar{L}$.

Remark 1 (the theorems agree at the boundary). $\alpha_J^* > 0$ if and only if $\varphi_J + C_O < \underline{p} + \pi_B$, which at the floor price is *exactly* $w > 0$: Theorem 2 is the $\alpha_J \rightarrow 0$ boundary case of Theorem 3. When the wedge is negative no structural verification is needed at all ($\alpha_J^* \leq 0$); when it is positive, the required accuracy rises continuously from zero as the wedge grows.

Three comparative statics deserve emphasis.

Proposition 5 (premium inversion). $\partial \alpha_J^* / \partial \underline{p} > 0$: every euro of price above the participation floor raises the accuracy that structural verification must attain. The bond logic of Proposition 3 is reversed at the screening margin: a premium enlarges the prize a mimic keeps with probability $1 - \alpha_J$, and recovers it only with probability α_J .

Each face of the price has classic ancestors of its own (Section 2: Guasch–Weiss on the prize side, Becker–Stigler on the bond side); what appears not to have met in one mechanism before is the pair with opposite signs. Jointly they discipline procurement design: *pay premia on contracts whose compliance problem is moral hazard by a known counterparty; pay the reservation price on contracts whose problem is choosing the counterparty, unless the screened attribute is verified nearly perfectly*. A sovereignty tender that combines a fat margin with weak control-verification is the worst configuration in the model: the margin finances the washing it invites.

Remark 2 (structural and behavioural audits are complements in screening). The joint-fake constraint in Theorem 3 binds *because* x_J is large: a mimic that expects to be caught on control has little to lose by faking operations too (Lemma 1’s complementarity, transposed to mimicry). So behavioural audits, useless for screening on their own (Theorem 2), earn a *supporting* role exactly when structural verification is strong, the opposite of the substitution pattern that per-attribute intuition suggests.

Remark 3 (transparency and accuracy are substitutes). $\partial\alpha_J^*/\partial\varphi_J < 0$: anything that makes cosmetic control-compliance expensive lowers the audit accuracy the buyer needs; beneficial-ownership registries, look-through disclosure of parents and control agreements, and liability for misdeclaration all raise the effective φ_J and $\bar{L}$. Corporate-transparency law is verification infrastructure. The same logic extends from corporate to *software* transparency: for the code-integrity attribute of licensed technology (Section 6), open source and reproducible builds are what raise α above zero. No buyer can audit a closed hyperscaler codebase and its update stream (Thompson 1984), so absent transparency-by-construction that attribute can only be secured structurally, by changing who supplies the code.

Remark 4 (the genuine-compliance route). With $C_J(B)$ finite, B has a third option: actually cede control. If it takes it, the buyer’s objective is *achieved*: a “mimic” that genuinely severs reach is success under another name. The no-washing constraints then only need to beat B ’s best *honest* option, $\max\{0, \underline{p} + \pi_B - C_J(B) - C_O\}$, which is easier the smaller the control gap $g_J = C_J(B) - C_J(G)$. Screening is hard precisely when g_J is enormous *and* $e > 0$, which is the sovereign-cloud configuration.

Proposition 6 (many mimic types: the worst type pins the accuracy). *Let potential mimics m differ in their ancillary rents π_m (hence stakes $\Omega_m = p + \pi_m + L$) and define each mimic’s wedge w_m as in Theorem 2. (a) Under behavioural-only verification, the washing set is exactly $\{m : w_m > 0\}$, independent of price and audits, type by type. (b) Under structural verification, the required accuracy is set by the **highest-stake** mimic: $\alpha_J \geq \max_m [1 - (\bar{L} + \varphi_J + C_O)/\Omega_m]$, increasing in the upper tail of the rent distribution. The screening problem is a worst-case problem: it does not matter how many well-behaved suppliers a market has; the accuracy bar is set by the single deepest-pocketed mimic: for sovereignty procurement, the largest hyperscaler’s ecosystem stake rather than the average bidder’s.*

5.3 Requirement design: structural-first

Assembling Parts I and II:

Corollary 2 (the optimal requirement set is sparse and structural). *An attribute contributes to screening only through the product of its type-cost gap and its verifiability; attributes with $g_i \approx 0$ (residency, operations) contribute zero at any accuracy (Theorem 2), while adding them to R raises the audit bill superadditively (Corollary 1), raises the mimic’s fake-gain menu, and, through the capacity constraint of Proposition 1(a), crowds the deterrence of the attributes that do screen. The buyer’s optimal tender requires and verifies effective control, holds behavioural requirements to those with direct payoff value net of their discipline cost, and concentrates the audit per Proposition 2.*

In institutional language: a sovereignty requirement must name the *control facts* (ultimate ownership, governing law, key custody, continuity under foreign compulsion) as eligibility conditions, and must verify them against the public record and the contract’s technical annexes, because those are the facts on which a sovereign and a non-sovereign supplier genuinely differ and which auditors can actually reach. Every other clause is either dead weight or camouflage.

Remark 5 (anti-screening requirements). The wedge accounting of Theorem 2 extends to any added requirement k : when both types comply genuinely, the mimicry wedge changes by $C_k(G) - C_k(B)$. A requirement that burdens the genuine type *more* than the mimic (fixed-cost compliance overhead, documentation, certification process, which a large incumbent’s compliance machinery absorbs more easily than a small genuine supplier) has a **wrong-signed gap**: it *widens* the wedge and actively selects against the objective. Bulky referentials are therefore worse than diluting (Corollary 1): on this margin they anti-screen, sorting by organizational size rather than by the screened property. The design rule sharpens accordingly: a scheme’s screening power lives in its structural core; its length is cost, and past the point where compliance overhead weighs more on the genuine type, negative screening.

5.4 The integrity pole: when the buyer audits only control

The baseline regime (A1)–(A3) assumes operational temptation is universal (even the genuine type would fake operations), so an operational discipline floor survives in every mechanism, and “audit only control” is impossible by construction. The opposite pole is equally coherent and, for sovereignty, arguably closer to the institutional reality: the type-gap extends to *every* attribute. Replace (A2)–(A3) with:

- (A2′) *Integrity.* $C_O(G) \leq \varphi_O < C_O(B)$: the genuine type performs operations authentically as a byproduct of being what it is (no temptation), while the non-sovereign type finds genuine operations-under-severed-control dearer than their appearance.

Under (A2′) the winner needs no discipline (a selected G complies genuinely everywhere at zero audit), so audits serve *only* to deter B ’s mimicry, whose two fake routes give the constraints $\Omega_B(1 - x_J) \leq K_1 \equiv \bar{L} + \varphi_J + C_O(B)$ and $\Omega_B(1 - x_J)(1 - x_O) \leq K_2 \equiv \bar{L} + \varphi_J + \varphi_O$, with $K_2 < K_1$. Define $\alpha_J^* = 1 - K_1/\Omega_B$ (as in Theorem 3, with $C_O(B)$ in place of C_O) and $\alpha_J^{**} = 1 - K_2/\Omega_B > \alpha_J^*$.

Theorem 4 (audit-only-control). Under (A1), (A2′), the polar control gap, $p = \underline{p}$ and $L = \bar{L}$:

- If $\alpha_J < \alpha_J^*$, no audit policy screens (washing, as in Theorem 2).
- If $\alpha_J \in [\alpha_J^*, \alpha_J^{**})$, screening is feasible but every no-washing mechanism uses behavioural audits ($x_O > 0$): control verification needs operational support, Remark 2 as a necessity.
- If $\alpha_J \geq \alpha_J^{**}$, the control-only screen ($x_J = \alpha_J^{**}$, $x_O = 0$) is feasible, and it is optimal iff $c_J K_1 \leq c_O \Omega_B$. Since $\Omega_B > K_1$ wherever screening is needed at all (the wedge is positive), $c_J \leq c_O$ suffices: whenever control is weakly cheaper to verify per unit of detection and its verification is accurate enough, the optimal mechanism audits only the control attribute.

Three readings. First, this is the “sufficient statistic” conjecture’s provable pole: with the type-gap on every attribute, the control audit, plus the *cost* wedge that the operational gap itself contributes to mimicry, screens everything, and operational audits are pure waste. Second, the integrity structure *shrinks* the mimicry wedge: B ’s cheapest mimicry now pays either the operational-authenticity gap $g_O = C_O(B) - C_O(G)$ (comply genuinely) or the joint-fake detection risk, so the effective wedge is $w' = e - (\varphi_J - C_J(G)) - g_O$. Screening is therefore easier at the integrity pole than under behavioural parity, and every unit of genuine operational advantage the sovereign sector holds does screening work for free. Third, the band $[\alpha_J^*, \alpha_J^{**})$ gives Remark 2 sharp boundaries: behavioural audits are *necessary* exactly when control verification is good enough to screen yet too weak to screen alone. Figure 3 displays the three regimes.

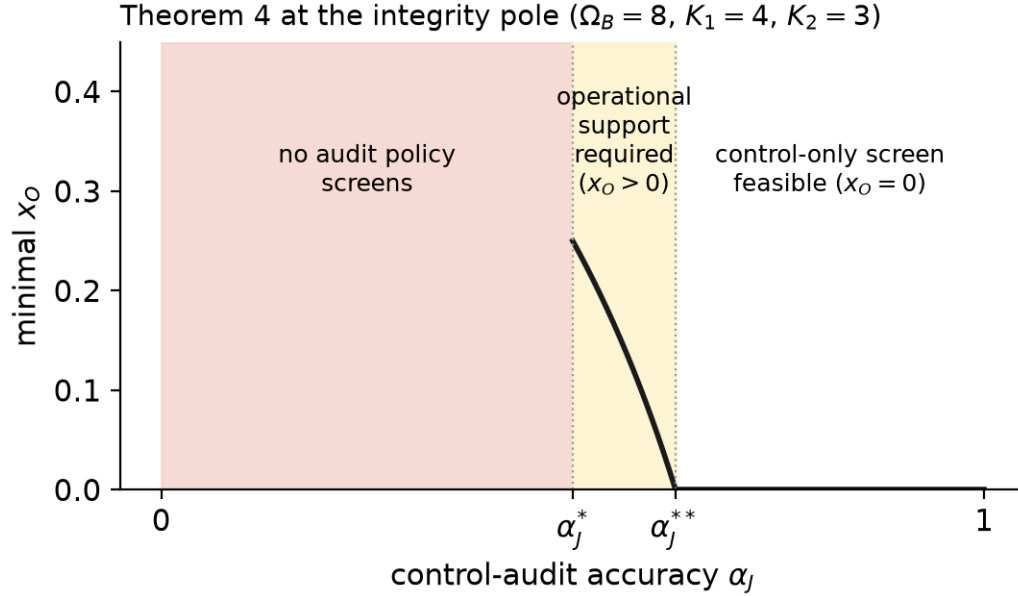


Figure 3. Theorem 4 at the integrity pole ($\Omega_B = 8, K_1 = 4, K_2 = 3$, so $\alpha_J^* = 0.5$ and $\alpha_J^{**} = 0.625$). The curve is the minimal behavioural-audit intensity x_O compatible with no washing. Below α_J^* no audit policy screens; on the band $[\alpha_J^*, \alpha_J^{**})$ every no-washing mechanism needs $x_O > 0$ from α_J^* on, the control-only screen is feasible.

The two poles now bracket the intermediate question, and Section 8 resolves it in both directions: with *independent* operational evidence, a second probe can be essential even when marginally dominated (Example 1: the discipline role is irreducible, as at the independence pole); when the operational audit merely *re-reads the control evidence with noise*, it is redundant (Theorem 5: the integrity pole's logic, stochasticized). The boundary is characterized in Section 8 (Proposition 7's essentiality sandwich; Remark 6's tightness); what remains is the application judgment of *where* evidence-sufficiency holds.

6. The public record

The model is a theory exercise; its discipline is internal (the failure regions of Propositions 1(c) and 4, and Theorem 2's wedge condition, are the ways it can come out against a sovereignty mandate). Still, its ordering predictions can be confronted with public facts, in the modest sense that the model *rationalizes* the documented record: no identification is claimed, the record is also consistent with simpler stories, and we say which facts would distinguish them.

Sorting. The verification regimes sort as predicted, on exactly the attributes they verify. France's SecNumCloud is a *security* qualification (ANSSI's visa for hosting sensitive data) rather than a sovereignty label, but its 3.2 revision embedded a page of structural criteria (non-EU ownership caps, no non-EU control, immunity from extra-European law, key custody), and for the sovereignty objective that page is the operative screen: $\alpha_J > 0$ pointed at the jurisdictional-control gap. It sorts at the top of the market: no hyperscaler-parented offering can pass it, while EU-controlled operators can and do (S3NS under full Thales ownership, December 2025; OVHcloud; Scaleway; Bleu, EU-owned and in the same licensed-technology class as S3NS, pending rather than excluded). The EU-level EUCS moved the other way: the March 2024 draft dropped its embedded sovereignty/immunity criteria, from a Theorem 3 regime to a Theorem 2 regime ($\alpha_J \rightarrow 0$), and under the resulting behavioural-only requirements the most operationally capable non-sovereign suppliers qualify, as Theorem 2 predicts. Two model mechanisms are visible inside the SecNumCloud record itself. First,

the *bulk effect*: the surrounding security referential (278 requirements; qualification costs and durations documented in the French Senate’s 2025 procurement inquiry, rapport n° 830) is legitimate for the security objective, where operational quality is itself the certified object, yet it prices small genuine suppliers out of demonstrating their sovereignty at all: Remark 5’s wrong-signed gap on top of Corollary 1’s dilution, and the case for a lean structural qualification unbundled from security tiers, the design the Cloud Sovereignty Framework’s separate sovereignty ranking begins to instantiate. Second, *scope discipline*: a screen delivers only the attributes it requires and verifies. S3NS severs jurisdictional control while operating licensed Google technology, which leaves two attributes SecNumCloud does not test: supply continuity (a foreign order to stop supporting the licensed stack starts a security-rot countdown) and code integrity, where verifying the absence of compelled covert access in a closed codebase of hundreds of millions of lines is beyond any licensee’s capacity (no full source access, no reproducible builds; Thompson 1984), so the attribute sits in Proposition 1(c)’s $\alpha \approx 0$ regime, exposed at any price and any audit budget. The qualification purchases *legal* access-immunity, which is real and more than any other European instrument delivers, and neither covert-access immunity nor technological autonomy; the two escapes are Remark 3’s (make the attribute verifiable by construction: open source, reproducible builds) and the structural one (a technology supplier beyond the compelling jurisdiction’s reach), which is what the tiered instruments’ top levels require (the Cloud Sovereignty Framework’s SEAL levels; the proposed Cloud and AI Development Act’s top tier). The institutional detail compressed here is developed in the companion policy brief.

The wedge on the record. Microsoft’s EU Data Boundary (completed February 2025) is full behavioural compliance, and the company’s own France public-affairs director testified to the French Senate (10 June 2025) that he could not guarantee French data would never be disclosed under a US CLOUD Act order: nominal compliance and effective gap, on the record. AWS’s European Sovereign Cloud (generally available January 2026) is the φ_J technology made flesh: EU incorporation, EU staff, and EU leadership under a US ultimate parent. Both are the configuration Proposition 1(c) and Theorem 2 say a location-and-practice test cannot screen out.

What would count against the model. A behavioural-only regime producing durable separation (residency requirements that non-sovereign suppliers systematically decline to satisfy); or a structural regime (SecNumCloud-like) whose qualified set includes suppliers under demonstrated foreign control. Neither appears in the record to date.

The one prediction the public record cannot reach is the counterintuitive one: raising behavioural-audit accuracy without touching control verification does not reduce washing (Theorem 2’s comparative static in α_O). No observed regime has moved α_O in isolation; we flag this as the discriminating test a future quasi-experiment would target.

7. Discussion: robustness and limitations

The bounded one-strike sanction. Every Part I result rides on the sanction being bounded and indivisible, and Section 4.1 asserted that the complementarity degrades continuously as per-count sanctioning capacity grows. The following one-parameter family displays the claim exactly.

Proposition 8 (graded sanctions). *Let the total sanction upon M detected fakes be $\sigma(M) = \tau \min(M, K)$: a per-count fine $\tau > 0$ capped at capacity $K \in \{1, 2, \dots\}$ counts; the baseline model is $(\tau, K) = (\Omega, 1)$. Full compliance on R is incentive-compatible iff $\tau \mathbb{E}[\min(M_D, K)] \geq \Delta(D)$ for every $D \subseteq R$, where M_D counts detections on the fake set D . At the per-attribute floors $x_j = \Delta_j/\tau$:*

- (i) *the deficit of the joint-deviation constraint is exactly $\tau \mathbb{E}[(M_D - K)^+]$, the expected number of detections beyond sanction capacity, priced at the fine;*
- (ii) *every deviation with $|D| \leq K$ is deterred by the per-attribute floors, so per-attribute analysis is*

exact for mandates of size at most K

(iii) the deficit is nonincreasing in K , strictly positive iff $|D| > K$ (with interior floors), and zero for $K \geq |R|$, where deterrence is additive and Corollary 1's superadditivity gap vanishes. At the baseline $(\Omega, 1)$ it equals $\Omega \mathbb{E}[(M_D - 1)^+] = V(D)$, Lemma 1's overlap value.

Complementarity, the cascade, and mandate overload therefore live exactly on the deviations larger than sanction capacity, and fade continuously as capacity grows. The procurement reality is $K = 1$: an award is voided once, debarment happens once, damages are capped by the firm's stake. A legal environment able to fine each violation separately at meaningful scale would restore per-attribute reasoning, and the proposition prices the capacity: a mandate of size n is safe for per-attribute analysis iff $n \leq K$. For $1 < K < n$ the binding constraints are the deviations of size exceeding K concentration persists in our numerical sweeps there, but the contra-polymatroid representation behind Theorem 1 is specific to $K = 1$ and we do not claim the cascade form beyond it.

Continuum of compliance and types. Largely discharged. Lemma 0 makes the binary compliance model without loss under detection-in-kind, and Proposition 6 extends the screening results to arbitrarily many mimic types (the worst type pins the accuracy). What remains open is the genuinely different problem where detection scales with the size of the gap: there the constraint family is a continuum and the cascade's contra-polymatroid structure has no off-the-shelf analogue (Conjecture 2).

Correlated attributes: the two poles and the middle. The paper now characterizes both poles of the type-attribute linkage. At the *independence* pole ((A2)–(A3): operational temptation universal, carrying no type information), operational audits retain an irreducible discipline floor and screening loads on control per Theorem 3. At the *integrity* pole ((A2'): temptation perfectly aligned with the type), the screen optimally collapses onto the single control audit (Theorem 4). The intermediate case, attributes stochastically linked to a one-dimensional type, is now resolved in both directions in Section 8: the naive version is *false* (an independent, marginally-dominated operational probe can be essential when deterrence saturates; Example 1), while the physical-garbling version is a *theorem* (Theorem 5: when the operational audit reads a degraded copy of the same evidence the control audit reads, dropping it is without loss). The boundary in between is now characterized as well, by the essentiality sandwich under independent probes (Proposition 7) and the tightness of evidence-sufficiency (Remark 6), and Şabac and Yoo (2018)'s warning is confirmed rather than contradicted: marginal informativeness ranking alone never licenses aggregation; a *sufficient* common evidence trail does.

Commitment. The buyer commits to audit probabilities; ex post it may be tempted to skip paid audits it expects to pass. Institutional commitment (certification schemes with mandated audit rates, third-party auditors under contract) approximates this; the no-commitment game and its refinements are future work (the model's separating logic does not rely on off-path beliefs precisely because of commitment).

Collusion and capture. We take (α_i, k_i) as technology. In reality the accuracy of a certification scheme is *chosen*, and lobbied. The companion paper endogenizes exactly this: the firms Theorem 3 would screen out lobby the standard-setter over which attributes are verified and how well; the EUCS reversal is the documented instance. Keeping capture out of the present paper is a deliberate scope choice.

False positives. Allowing $\alpha_i^{FP} > 0$ would penalize genuine compliance and create a countervailing force against auditing; with public-record structural facts the false-positive rate is plausibly tiny, which further favours loading on J the extension matters for the operational attribute, though, and belongs in the general model.

One supplier. Competition among non-sovereign incumbents dissipates π_B toward zero and *relaxes* the screening problem (Theorem 3's threshold falls with π_B); competition from the sovereign entrant disciplines prices (Proposition 4). Modeling the full auction is future work; the one-firm case is the conservative one for our claims.

What would kill the paper's central claim. If, in the general model, the optimal mechanism turned out to secure behavioural attributes as cheaply as structural ones (no concentration), or if screening survived $g_J \rightarrow \varepsilon$ (washing possible only by assumption), the structural-first message would be an artifact. The binary model says otherwise for transparent reasons; the generalizations in Section 8 are where that verdict could still be overturned.

8. When is one audit enough?

The two poles of Section 5.4 bracket an intermediate question: when operational evidence is stochastically linked to the type, can the screening audit ignore it? This section resolves the question in both directions and states what remains open.

Example 1 (the weak sufficient-statistic conjecture is false). An earlier draft conjectured: if, conditional on every type–action profile, the operational audit's signal is a Blackwell garbling of the control audit's signal, and $k_J \leq k_O$, then every optimal mechanism sets $a_O = 0$. Two findings kill it. *First*, the profile-wise premise hides a structural consequence: across two profiles differing only in the operational act, if the control audit's distribution is unchanged (a control inspection does not see operational shortcuts), the fixed-kernel garbling equation forces the operational audit's distribution to be unchanged too: under the weak premise the O-audit is **act-blind wherever the J-audit is**, so it can only ever serve as a second type-probe. *Second*, even as a type-probe the conjecture fails: a probe Blackwell-dominated *marginally* but drawn *independently* given the type adds information no reading of the J-evidence replicates. In a verified numerical instance, the mimic's prize is large enough that deterrence **saturates** at the control audit's maximal accuracy (no J-only mechanism screens at any price and any policy, by a mechanism-free bound) while running both audits deters: the dominated independent probe is *essential*. This is Şabac and Yoo (2018) materializing in our setting: a marginal informativeness ranking does not license aggregation.

Theorem 5 (audit redundancy under physical garbling). *Model the audits as readers of a common latent body of evidence: $z \sim F(\cdot \mid \text{type, acts})$ (for sovereignty, the control facts and their traces); the control audit reads $s_J = z$ at cost k_J the operational audit reads a degraded copy $s_O = h(z, \varepsilon)$, ε exogenous independent noise, at cost $k_O \geq k_J$. Then every mechanism in the class of §3.1, with report- and signal-contingent, sequential audit policies, is weakly dominated by one that never runs the operational audit.* The proof (Appendix) couples by simulation: replace each O-run by a J-run plus a synthetic draw $\tilde{s}_O = h(s_J, \hat{\varepsilon})$ the joint law of all signals is identical profile-by-profile, so **every constraint is unchanged rather than merely relaxed**, and audit costs weakly fall. This is the “master-key” reading of sovereignty made formal (everything an operational inspection can find is a noisy shadow of the control facts), and Theorem 4 is its deterministic special case. A fuller derivation, with flags for the reviewer, is available from the author as a companion note.

Proposition 7 (independent probes: the essentiality sandwich). *In the two-probe screening environment (flags independent given the type; $k_J \leq k_O$), define the saturation bound at the no-false-positive floor price $\underline{p}_0 = \max\{0, C_G - \pi_G\}$:*

$$E = (1 - d_J)(\underline{p}_0 + \pi_B) - d_J \bar{L} - \varphi_J - C_O,$$

and let E_1 be the same expression at the full-intensity floor price $\underline{p}_1 = (C_G + f_J \bar{L})/(1 - f_J) - \pi_G$. Then: (a) if $E > 0$, no mechanism using only the control audit screens, by a mechanism-free bound, so the operational probe is essential;* (b) *if $E_1 \leq 0$, the full-intensity control-only mechanism screens;*

the band $E \leq 0 < E_1$, feasibility is instance-specific (audit intensity trades detection against false-positive price inflation). (c) The essentiality region expands in the mimic's rents π_B and contracts in d_J , $\bar{L}$, φ_J , and C_O .^{*} With binary flags, the interior use-the-second-probe-or-not comparison reduces to a finite enumeration over the portfolio's ROC vertices (AND, single-audit, OR, and mixtures); we omit the taxonomy. Figure 4 displays the sandwich.

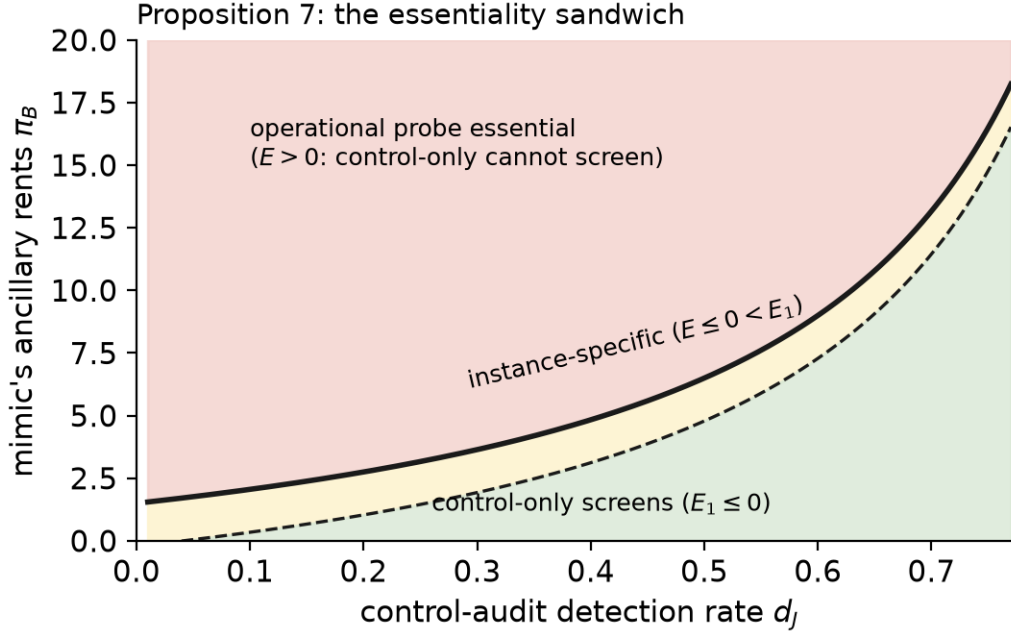


Figure 4. The essentiality sandwich (Proposition 7) in the (d_J, π_B) plane ($\bar{L} = 3$, $\varphi_J = C_O = 1$, $C_G = 1$, $\pi_G = 0.5$, $f_J = 0.3$). Above the solid curve ($E > 0$) no control-only mechanism screens and the operational probe is essential; below the dashed curve ($E_1 \leq 0$) the full-intensity control-only mechanism screens; between the curves feasibility is instance-specific. The essential region expands with the mimic's rents and contracts with the control audit's detection rate.

Remark 6 (Theorem 5 is tight, and slightly stronger). Theorem 5's proof uses only that the control audit's signal is *sufficient* for the latent evidence ($s_O \perp \text{profile} \mid s_J$): under sufficiency the simulation kernel is profile-free and the coupling runs verbatim. And that condition is exact: if the control audit is itself a **noisy reader** ($s_J = m(z, \eta)$, no longer sufficient), the two audits become conditionally independent readers of z , Proposition 7 applies with the readers' operating characteristics, and in the saturation region the operational audit is again essential, as verified in a noisy-reader instance where control-only screening is infeasible while the two-audit portfolio screens. **No informativeness-per-euro ordering of noisy readers yields a one-audit theorem; only evidence-sufficiency does.**

What remains open here is an application judgment, the mathematics being settled: whether evidence-sufficiency holds for sovereignty *attribute-by-attribute*, plausibly yes for access-related operations (the master-key reading) and plausibly no for the update-stream/code-integrity channel (Section 6). Attribute-by-attribute sufficiency would yield a *partitioned* Theorem 5: audit only control where the operational trail is its shadow; audit operations where the trail is independent evidence.

Conjecture 2 (gap-dependent detection). With continuous compliance where detection probability increases in the size of the shortfall (plausible for operational attributes), the incentive family becomes a continuum of constraints of the form “additive savings $\leq$ convex transform of the aggregate shortfall,” and we conjecture the cost-minimizing audit remains a cascade along a priority

order, with Lemma 0's all-or-nothing property failing gracefully (fakes concentrate on a shrinking set of attributes as detection curvature rises). This is a semi-infinite program; the contra-polymatroid machinery does not apply off the shelf.

9. Conclusion

A requirement is a screen only if it is costly to counterfeit and cheap to verify, and in multidimensional compliance problems those two properties interact in ways per-attribute intuition misses: fakes are complements, deterrence is a common resource, audits should concentrate, the price disciplines and corrupts on different margins, and no accuracy pointed at the wrong attribute can substitute for any accuracy pointed at the right one. Applied to Europe's sovereignty procurement, the model gives the debate a precise vocabulary: data residency and operational requirements are non-screens rather than weak screens (Theorem 2); incorporation without severed control is the cosmetic technology itself (Proposition 1(c)); the integrity of licensed closed code is unverifiable at any audit budget, so it must be secured structurally or made verifiable by construction, through open source and reproducible builds (Proposition 1(c) again, Remark 3); verification of effective control is feasible, and its required accuracy falls with corporate transparency and rises with the margins on offer (Theorem 3, Proposition 5); mandates whose verification cost exceeds their exposure value should not be issued (Proposition 4). The model was built so that it could recommend against sovereignty mandates, and in a characterized region it does. Where it recommends them, it says exactly what they must verify.

Appendix A. Proofs

Throughout, $\Omega = p + \pi + L$, $x_i = a_i \alpha_i$, $S(D) = \prod_{j \in D} (1 - x_j)$, $\delta_i = \Delta_i / \Omega$, $c_i = k_i / \alpha_i$, $\bar{\gamma} = 1 - \delta_J - \delta_O$. All closed-form claims are additionally verified against independent brute-force computation; the verification suite accompanies the paper.

Proof of Remark 0. (a) Fix such a mechanism and an equilibrium in which some participation strategy σ_0 , followed by losing, carries net expected transfer $t > \kappa$. A non-sovereign entrant can execute σ_0 's observable actions exactly: reports and fee payments are unrestricted, and a non-winner delivers no service and faces no audit (audits attach to the delivered contract; §3.1 timing), so no instrument of the mechanism ever engages the entrant's type. Each additional entrant playing σ_0 therefore nets at least $t - \kappa > 0$ in expectation, and the buyer's expected outlay on non-winners grows at least linearly in the number of entrants, violating any finite budget. If instead the mechanism's rules cap total non-winner payments, entry continues until the per-entrant expected transfer falls to κ , with the same conclusion: in free-entry equilibrium, no non-winning strategy carries rents. (b) Consider any scheme with an application fee F and a rebate $\rho \leq F$ to non-winners, both measurable in reports and the award outcome. In the coupling of Theorem 2(a), B 's shadow strategy makes the identical reports and the identical fee payments as G 's strategy; the distribution of awards, and hence of rebates, is identical; so the fee-and-rebate terms cancel from the payoff difference, which remains exactly w . An instrument that charges both types identically on identical observables cannot enter the wedge. $\square$

Proof of Lemma 0. Fix any audit profile and any claim vector, and let $D = \{i : \text{shortfall}_i > 0\}$. Under detection-in-kind the detection distribution (hence the probability of voiding and the expected sanction) depends on the shortfall vector only through D . Conditional on D , the firm's payoff is (expected award terms, fixed) plus savings strictly increasing in each shortfall on D so every best response takes each shortfall on D to its maximum and each off D to zero. $\square$

Proof of Lemma 1. $u(D \cup \{j\}) - u(D) = \Omega S(D)(1 - x_j) - \Omega S(D) + \Delta_j = \Delta_j - \Omega x_j S(D)$. $S(D)$ is decreasing in D , so the difference is increasing in D (increasing differences; u is supermodular on 2^R). For the second claim it suffices that $1 - \prod_{j \in D} (1 - \delta_j) < \sum_{j \in D} \delta_j$ whenever $|D| \geq 2$ and all $\delta_j \in (0, 1)$. Induction on $|D|$: for $D = \{i, j\}$, $1 - (1 - \delta_i)(1 - \delta_j) = \delta_i + \delta_j - \delta_i \delta_j < \delta_i + \delta_j$ and if the claim holds for D then $1 - \prod_D (1 - \delta_j)(1 - \delta_c) = [1 - \prod_D (1 - \delta_j)] + \delta_c \prod_D (1 - \delta_j) < \sum_D \delta_j + \delta_c$ since $\prod_D (1 - \delta_j) < 1$. Multiplying by Ω : $\Omega[1 - \prod(1 - \delta_j)] < \Delta(D)$, i.e. (IC- D) fails strictly at the floors. $\square$

Proof of Proposition 1. (a) *Necessity*: $x_j \leq \alpha_j$ caps $1 - S(D)$ at $1 - \prod_{j \in D} (1 - \alpha_j)$ if the stated inequality fails for some D , (IC- D) fails at every feasible audit profile. *Sufficiency*: at $a_j = 1$ ($x_j = \alpha_j$) every (IC- D) holds by assumption; participation holds for $p \geq \sum_{i \in R} C_i - \pi$. (b) If all $\alpha_i > 0$ then $1 - \prod_{j \in D} (1 - \alpha_j) > 0$ for every nonempty D , and the right side of (a) grows linearly in p while the left is constant: all constraints hold for p large. If some required i has $\alpha_i = 0 < \Delta_i$, take $D = \{i\}$: the right side is $0 < \Delta_i$ at every p . (c) With $\alpha_i = 0$, faking i is undetectable, so for any strategy profile, switching i from genuine to fake raises the firm's payoff by $\Delta_i > 0$: every best response fakes i the firm participates whenever the mechanism is acceptable at all, and fakes i . $\square$

Proof of Proposition 2. If the joint constraint were slack at an optimum with some x_i above the floors, reducing that x_i preserves feasibility and lowers cost; so either both x_i sit at floors (infeasible by Lemma 1) or the joint constraint binds. On the binding curve, substitute $x_O = 1 - \bar{\gamma}/(1 - x_J)$: the cost $c_J x_J + c_O[1 - \bar{\gamma}/(1 - x_J)]$ has second derivative $-2c_O \bar{\gamma}/(1 - x_J)^3 < 0$, strictly concave, so the minimum over the feasible interval of x_J is at an endpoint. The endpoints are $x_O = \delta_O$ (whence $x_J = 1 - \bar{\gamma}/(1 - \delta_O) = \delta_J/(1 - \delta_O)$, using $(1 - \delta_O) - \bar{\gamma} = \delta_J$) and symmetrically $x_J = \delta_J$, or an accuracy cap where applicable. Corner costs: $\text{Cost}_A = c_J \delta_J/(1 - \delta_O) + c_O \delta_O$, $\text{Cost}_B = c_J \delta_J + c_O \delta_O/(1 - \delta_J)$ their difference is $\delta_J \delta_O [c_J/(1 - \delta_O) - c_O/(1 - \delta_J)]$, negative iff $c_J(1 - \delta_J) < c_O(1 - \delta_O)$, i.e. $\eta_J < \eta_O$. Substituting δ 's gives (4.1). Feasibility of the loaded corner requires $\delta_J/(1 - \delta_O) \leq \alpha_J$ otherwise $x_J = \alpha_J$ and the joint constraint pins $x_O = 1 - \bar{\gamma}/(1 - \alpha_J)$. $\square$

Proof of Corollary 1. From (4.1), $A = c_J \Delta_J/(\Omega - \Delta_O) + c_O \Delta_O/\Omega > c_J \Delta_J/\Omega + c_O \Delta_O/\Omega$, the standalone (single-requirement) audit costs at the same stake, since $\Omega - \Delta_O < \Omega$. Adding a requirement with $\Delta > 0$ weakly raises every existing floor's inflation factor and strictly raises the loaded attribute's. $\square$

Proof of Theorem 1. Change variables $z_i = -\ln(1 - x_i) \in [0, \infty)$. Then (IC- D) reads $z(D) \equiv \sum_{j \in D} z_j \geq \zeta(D) \equiv -\ln(1 - \Delta(D)/\Omega)$ (finite under implementability with headroom), and the audit cost is $\sum_i c_i(1 - e^{-z_i})$: separable, strictly increasing, strictly concave. Now $\zeta(D) = f(\Delta(D))$ with $f(t) = -\ln(1 - t/\Omega)$ strictly convex and increasing and $\Delta(\cdot)$ additive; a convex increasing transform of an additive set function has increasing differences, so ζ is supermodular, monotone, with $\zeta(\emptyset) = 0$. The feasible region $P = \{z \geq 0 : z(D) \geq \zeta(D) \forall D\}$ is therefore a contra-polymatroid; since $\zeta(\{i\}) > 0$ for all i , the nonnegativity constraints never bind, and the extreme points of P 's lower boundary are exactly the chain-greedy vectors z^σ with $z_{\sigma(k)}^\sigma = \zeta(S_k) - \zeta(S_{k-1})$ over orderings σ (Edmonds 1970; Fujishige 2005), each feasible for the full family by supermodularity. The objective is bounded, increasing, and concave, so its minimum over P is attained at such a vertex. Transforming back, $1 - x_{\sigma(k)} = e^{-z_{\sigma(k)}^\sigma} = \frac{1 - \Delta(S_k)/\Omega}{1 - \Delta(S_{k-1})/\Omega}$, i.e. $x_{\sigma(k)} = \Delta_{\sigma(k)}/(\Omega - \Delta(S_{k-1}))$, the cascade. For the ordering: exchanging adjacent attributes $a = \sigma(k)$, $b = \sigma(k+1)$ at prefix $S = S_{k-1}$ changes only their two terms, and

$$\text{Cost}(a b) - \text{Cost}(b a) = \frac{\Delta_a \Delta_b}{\Omega - \Delta(S)} \left[\frac{c_b}{\Omega - \Delta(S) - \Delta_a} - \frac{c_a}{\Omega - \Delta(S) - \Delta_b} \right],$$

so a -before- b is (weakly) better iff $c_a (\Omega - \Delta(S) - \Delta_a) \geq c_b (\Omega - \Delta(S) - \Delta_b)$: the larger prefix-adjusted index goes first (equivalently, the smaller one is audited later, harder). The comparison depends on $\Delta(S)$, so no prefix-free index orders the cascade in general; at $n = 2$, $S = \emptyset$ and it reduces to $\eta_a \geq \eta_b$, recovering Proposition 2 exactly (the low- η attribute is loaded). $\square$

Proof of Proposition 3. L enters payoffs only through Ω and deterrence is increasing in Ω at zero marginal cost, so $L = \bar{L}$. $A(p)$ from (4.1) is strictly decreasing and strictly convex in p (term-by-term inverse-linear), so $T(p) = p + A(p)$ is strictly convex; $p^* > \underline{p}$ iff $T'(\underline{p}) < 0$, which is the displayed condition; the FOC is $1 = -A'(p^*) = c_J \Delta_J / (\Omega^* - \Delta_O)^2 + c_O \Delta_O / (\Omega^*)^2$ with $\Omega^* = p^* + \pi + \bar{L}$. (Regime: under $c_J \leq c_O$ and $\Delta_J \geq \Delta_O$, both factors of $\eta_i = c_i(1 - \Delta_i/\Omega)$ are ranked in J 's favour for every feasible Ω , so $\eta_J \leq \eta_O$ at every price and Proposition 2 loads J throughout; since the cascade intensities fall as Ω rises, interior feasibility at $\underline{p}$ implies interior feasibility along the whole path, and $A(p)$ is given by (4.1) at every price considered.) $\square$

Proof of Proposition 4. Direct comparison of the three welfare expressions, with A from (4.1); the inclusion of the joint-deviation burden makes A strictly larger than the per-attribute sum (Corollary 1), enlarging both stated regions. $\square$

Proof of Theorem 2. (a) Fix any procurement mechanism and any strategy σ available to G (a choice of contract from the menu, fee payments, reports, and genuine compliance $s_J = s_O = 1$). Let B play the *shadow strategy* $\hat{\sigma}$: identical choices, genuine compliance on O , cosmetic compliance on J . Every audit outcome is distributed identically under (σ, G) and $(\hat{\sigma}, B)$: J -audits detect nothing ($\alpha_J = 0$), O -audits face genuine compliance under both. Hence every transfer, award probability, and penalty event is identical, and the payoff difference is the cost-and-rents difference only: $u_B(\hat{\sigma}) - u_G(\sigma) = (\pi_B - \pi_G) + C_J(G) - \varphi_J + (C_O - C_O) = w$. (b) If $w > 0$: in any equilibrium in which some strategy giving G a nonnegative payoff leads, with positive probability, to an award requiring $\hat{s}_J = 1$, the type- B firm obtains $w > 0$ by shadowing it, so it does; the mechanism cannot condition on the type (identical observables), so the award reaches B with $\hat{s}_J = 1 > s_J = 0$ with positive probability: washing. The only equilibria without washing are those in which no sovereignty-requiring award is ever made. (c) If $w \leq 0$: post the single contract $(\underline{p}, x_O = \Delta_O/\Omega_G \leq \alpha_O)$ by (A3), $L = \bar{L}$). Then $u_G = 0$ B 's fake- J deviation gives $u_G + w \leq 0$ faking O as well changes the payoff by $\Delta_O - \Omega_B x_O(1 - x_J)|_{x_J=0} = \Delta_O - \Omega_B x_O \leq \Delta_O - \Omega_G x_O \leq 0$, so the joint fake is dominated. No washing. $\square$

Remark (Theorem 2 boundary). The condition $\alpha_J^* > 0$ of Theorem 3 reads $\varphi_J + C_O < \underline{p} + \pi_B$ substituting $\underline{p} = C_J(G) + C_O - \pi_G$ gives $\varphi_J - C_J(G) < e$, i.e. $w > 0$: the two theorems agree at $\alpha_J \rightarrow 0$.

Proof of Theorem 3 and Proposition 5. No-washing requires $u_B(\{J\}) = \Omega_B(1 - x_J) - L - \varphi_J - C_O \leq 0$ with $\Omega_B = p + \pi_B + L$, i.e. $x_J \geq 1 - (L + \varphi_J + C_O)/\Omega_B$, and the joint-fake constraint $\Omega_B(1 - x_J)(1 - x_O) \leq L + \varphi_J + \varphi_O$. Unlike the $\alpha_J = 0$ case of Theorem 2(c), the joint fake is *not* dominated by the single fake when $x_J > 0$: $u_B(\{J, O\}) - u_B(\{J\}) = \Delta_O - \Omega_B(1 - x_J)x_O$, which turns positive at large x_J : Lemma 1's complementarity, now on the mimicry margin (a mimic likely to be caught on J fakes O almost for free). Both constraints are therefore imposed; each is relaxed by raising x_J , x_O , or L and tightened by raising p , so a no-washing mechanism exists iff both hold at $x_J = \alpha_J$, $x_O = \alpha_O$, $L = \bar{L}$, $p = \underline{p}$, which is the display. Comparative statics of the threshold $\alpha_J^*(p) = 1 - (L + \kappa)/\Omega_B$ with $\kappa = \varphi_J + C_O$: increasing in Ω_B 's components p and π_B (numerator fixed; this is Proposition 5) and decreasing in φ_J in L , $\partial \alpha_J^* / \partial L = [\kappa - p - \pi_B] / \Omega_B^2$, negative exactly in the screening-relevant region $\kappa < p + \pi_B$ (Remark 1), so a higher penalty cap lowers the threshold. $\square$

Proof of Corollary 2. Immediate from Theorem 2 (zero screening contribution of $g_i = 0$ attributes at any α), Corollary 1 (superadditive audit cost of additional requirements), and Proposition 1(a) (common deterrence capacity). $\square$

Proof of Proposition 6. (a) Theorem 2(a)'s shadow-strategy argument uses only mimic m 's own costs and rents, so it holds type by type: $u_m(\hat{\sigma}) - u_G(\sigma) = w_m$ in every procurement mechanism with $\alpha_J = 0$ mimic m enters iff $w_m > 0$. (b) With $\alpha_J > 0$, mimic m 's fake- J constraint is $\Omega_m(1 - x_J) \leq \bar{L} + \varphi_J + C_O$, whose right side is common and whose left side is increasing in Ω_m the system over all entering mimics is therefore equivalent to the single constraint at $\max_m \Omega_m$. $\square$

Proof of Theorem 4. Under (A2'), $\Delta_i(G) \leq 0$ on every attribute, so a selected G 's best response is genuine compliance at any audit profile: no discipline constraints. B 's mimicry deviations (with the polar control gap, every mimicry fakes J the finite- $C_J(B)$ genuine- J route is priced out as in Remark 4) yield constraints (1) $\Omega_B(1 - x_J) \leq K_1$ and (2) $\Omega_B(1 - x_J)(1 - x_O) \leq K_2$. (i) Constraint (1) cannot be met for $x_J \leq \alpha_J < \alpha_J^*$, and x_O does not enter it. (ii) For $\alpha_J < \alpha_J^{**}$, at $x_O = 0$ constraint (2) requires $x_J \geq 1 - K_2/\Omega_B = \alpha_J^{**} > \alpha_J$: infeasible, so any no-washing mechanism has $x_O > 0$ feasibility itself holds since (2) can be met by raising x_O . (iii) For $\alpha_J \geq \alpha_J^{**}$, minimize $c_J x_J + c_O x_O$ subject to (1), (2): along the binding (2)-hyperbola the cost is strictly concave in x_J (as in Proposition 2), so the optimum is a corner: either $(x_J, x_O) = (1 - K_2/\Omega_B, 0)$ (control-only) or $(1 - K_1/\Omega_B, 1 - K_2/K_1)$ (constraint (1) binding, operations supporting). The cost difference is $(K_1 - K_2)[c_J/\Omega_B - c_O/K_1]$, so control-only is optimal iff $c_J K_1 \leq c_O \Omega_B$. Screening is needed only where B 's unaudited mimicry is profitable at the floor price, i.e. $\Omega_B > K_1$ there $c_J \leq c_O$ implies $c_J K_1 < c_O \Omega_B$. $\square$

Proof of Proposition 8. Detection events are independent Bernoulli(x_j) across $j \in D$, so $\mathbb{E}[M_D] = \sum_{j \in D} x_j$ and the gain from faking D is $\Delta(D) - \tau \mathbb{E}[\min(M_D, K)]$, giving the stated IC family. Write $\min(M, K) = M - (M - K)^+$. At the floors $x_j = \Delta_j/\tau$, $\tau \mathbb{E}[M_D] = \Delta(D)$, so the IC deficit is $\Delta(D) - \tau \mathbb{E}[M_D] + \tau \mathbb{E}[(M_D - K)^+] = \tau \mathbb{E}[(M_D - K)^+]$: (i). If $|D| \leq K$ then $M_D \leq K$ almost surely, the correction term vanishes, and (IC- D) holds with equality at the floors: (ii). Pointwise, $(M - K)^+$ is nonincreasing in K and zero for $K \geq |D|$ its expectation is strictly positive iff $\Pr(M_D > K) > 0$, i.e. iff $|D| > K$ and the floors are interior. At $(\tau, K) = (\Omega, 1)$, $\mathbb{E}[(M_D - 1)^+] = \mathbb{E}[M_D] - \Pr(M_D \geq 1) = \sum_{j \in D} \delta_j - 1 + \prod_{j \in D} (1 - \delta_j) = V(D)/\Omega$: (iii). $\square$

Proof of Theorem 5. Fix any mechanism M in the class of §3.1, with any audit policy (report- and signal-contingent, sequential) and outcome rule ψ measurable in reports and observed signals. Define M' : on every history where M runs the operational audit, M' runs the control audit instead if it has not already run, draws a synthetic $\tilde{s}_O = h(s_J, \hat{\varepsilon})$ with fresh independent noise, and applies M 's policy and rule verbatim, treating $\tilde{s}_O$ as the operational signal, including any continuation decisions of M that condition on it, in whatever order M observes signals. Under SG the true signals satisfy $(s_J, s_O) = (z, h(z, \varepsilon))$ with ε independent of everything given the profile; M' generates $(z, h(z, \hat{\varepsilon}))$ on the same events. For every type-action profile and every firm strategy, the joint law of all signals entering ψ is therefore identical under M and M' , hence so is the distribution of awards, transfers, and penalties for every player. Every participation and incentive constraint is preserved exactly (equality of signal laws, so more than a relaxation), equilibrium behaviour is unchanged, and the buyer's payoff differs only by the audit-cost saving: $k_O - k_J \geq 0$ per substituted audit, k_O per redundant one. Setting $a_O \equiv 0$ is without loss, and strictly improving whenever M ran the operational audit with positive probability and $k_J < k_O$ (or ran both). The load-bearing scope condition: the control audit reads the latent evidence exactly ($s_J = z$); if it is itself a noisy reader, the coupling fails and the environment falls to Proposition 7's regime (Remark 6). $\square$

Proof of Proposition 7. (a) Under any control-only policy the mimic's detection probability satisfies $D_B \leq a_J d_J \leq d_J$, and any mechanism the genuine type accepts pays $p \geq \underline{p}_0$. The mimic's payoff $(1 - D_B)(p + \pi_B) - D_B \bar{L} - \varphi_J - C_O$ is strictly decreasing in D_B (derivative $-(p + \pi_B) - \bar{L}$) and strictly increasing in p (derivative $1 - D_B$), so it is bounded below by its value at $(D_B, p) = (d_J, \underline{p}_0)$, which is $E > 0$: the mimic strictly profits under *every* control-only mechanism the genuine type accepts; washing, mechanism-free. (b) The full-intensity control-only mechanism ($a_J = 1$, trigger on the flag) has $D_G = f_J$, floor price $\underline{p}_1$, and mimic payoff exactly $E_1 \leq 0$: it screens. (c) $\partial E / \partial \pi_B = 1 - d_J > 0$ $\partial E / \partial d_J = -(\underline{p}_0 + \pi_B) - \bar{L} < 0$ the remaining derivatives are immediate. $\square$

Proof of Remark 6. *Sufficiency:* if $s_O \perp$ profile $| s_J$, the kernel $\mathcal{L}(s_O | s_J)$ is profile-free and known to the designer; replacing each O-run by a draw from it reproduces the joint law of all signals profile-by-profile, and the proof of Theorem 5 applies verbatim. *Tightness:* with $s_J = m(z, \eta)$ and $s_O = h(z, \varepsilon)$, noises independent given z , the two audits are conditionally independent informative readers of z taking z binary with type-dependent law maps the environment into Proposition 7's, with (d_i, f_i) the readers' operating characteristics, and any instance in the saturation region makes the operational audit essential. Verified instance: $\Pr(z = 1 | B) = 0.9$, $\Pr(z = 1 | G) = 0.05$, reader noises $(\eta_J, \eta_O) = (0.25, 0.30)$, $\pi_B = 14$: control-only screening is infeasible while the two-audit portfolio screens. $\square$

References

References were verified against Crossref, RePEc, arXiv, and publisher records; the verification log is available from the author. Reading status of the closest comparators is disclosed in Section 2.

- Amershi, A. H. & Hughes, J. S. (1989). Multiple signals, statistical sufficiency, and Pareto orderings of best agency contracts. *RAND Journal of Economics* 20(1):102–112. <https://doi.org/10.2307/2555654>
- Asseyer, A. & Weksler, R. (2024). Certification design with common values. *Econometrica* 92(3):651–686. <https://doi.org/10.3982/ECTA21653>
- Avenhaus, R., von Stengel, B. & Zamir, S. (2002). Inspection games. In *Handbook of Game Theory with Economic Applications*, vol. 3, ch. 51, Elsevier.
- Ball, I. (2025). Scoring strategic agents. *AEJ: Microeconomics* 17(1):97–129. <https://doi.org/10.1257/mic.20230275>
- Ball, I. & Kattwinkel, D. (2025). Probabilistic verification in mechanism design. *Theoretical Economics* 20(4):1247–1284. <https://doi.org/10.3982/TE6266>
- Becker, G. S. (1968). Crime and punishment: an economic approach. *Journal of Political Economy* 76(2):169–217. <https://doi.org/10.1086/259394>
- Becker, G. S. & Stigler, G. J. (1974). Law enforcement, malfeasance, and compensation of enforcers. *Journal of Legal Studies* 3(1):1–18. <https://doi.org/10.1086/467507>
- Ben-Porath, E., Dekel, E. & Lipman, B. L. (2014). Optimal allocation with costly verification. *American Economic Review* 104(12):3779–3813. <https://doi.org/10.1257/aer.104.12.3779>
- Bester, H. (1985). Screening vs. rationing in credit markets with imperfect information. *American Economic Review* 75(4):850–855.
- Bond, P. & Gomes, A. (2009). Multitask principal–agent problems: optimal contracts, fragility, and effort misallocation. *Journal of Economic Theory* 144(1):175–211.
- Calveras, A., Ganuza, J.-J. & Hauk, E. (2004). Wild bids: gambling for resurrection in procurement contracts. *Journal of Regulatory Economics* 26(1):41–68.
- Calzolari, G. & Spagnolo, G. (2009). Relational contracts and competitive screening. CEPR Discussion Paper 7434.

- Border, K. C. & Sobel, J. (1987). Samurai accountant: a theory of auditing and plunder. *Review of Economic Studies* 54(4):525–540. <https://doi.org/10.2307/2297481>
- Crocker, K. J. & Morgan, J. (1998). Is honesty the best policy? Curtailing insurance fraud through optimal incentive contracts. *Journal of Political Economy* 106(2):355–375. <https://doi.org/10.1086/250012>
- Dionne, G., Giuliano, F. & Picard, P. (2009). Optimal auditing with scoring: theory and application to insurance fraud. *Management Science* 55(1):58–70. <https://doi.org/10.1287/mnsc.1080.0905>
- Edmonds, J. (1970). Submodular functions, matroids, and certain polyhedra. In *Combinatorial Structures and Their Applications*, Gordon and Breach, 69–87.
- Eeckhout, J., Persico, N. & Todd, P. E. (2010). A theory of optimal random crackdowns. *American Economic Review* 100(3):1104–1135.
- Farrell, H. & Newman, A. L. (2019). Weaponized interdependence: how global economic networks shape state coercion. *International Security* 44(1):42–79. https://doi.org/10.1162/isec_a_00351
- Fujishige, S. (2005). *Submodular Functions and Optimization*, 2nd ed. Elsevier (Annals of Discrete Mathematics 58).
- Gale, D. & Hellwig, M. (1985). Incentive-compatible debt contracts: the one-period problem. *Review of Economic Studies* 52(4):647–663. <https://doi.org/10.2307/2297737>
- Georgiadis, G. & Szentes, B. (2020). Optimal monitoring design. *Econometrica* 88(5):2075–2107. <https://doi.org/10.3982/ECTA16475>
- Grossman, G. M. & Helpman, E. (1994). Protection for sale. *American Economic Review* 84(4):833–850. <https://www.jstor.org/stable/2118033>
- Grossman, G. M. & Shapiro, C. (1988). Counterfeit-product trade. *American Economic Review* 78(1):59–75.
- Guasch, J. L. & Weiss, A. (1981). Self-selection in the labor market. *American Economic Review* 71(3):275–284.
- Holmström, B. (1979). Moral hazard and observability. *Bell Journal of Economics* 10(1):74–91. <https://doi.org/10.2307/3003320>
- Huang, C. P.-C. (2026). Screening workers with affirmative action. Working paper. <https://arxiv.org/abs/2604.00615>
- Klein, B. & Leffler, K. B. (1981). The role of market forces in assuring contractual performance. *Journal of Political Economy* 89(4):615–641. <https://doi.org/10.1086/260996>
- Kuhn, M. & Siciliani, L. (2009). Performance indicators for quality with costly falsification. *Journal of Economics & Management Strategy* 18(4):1137–1154. <https://doi.org/10.1111/j.1530-9134.2009.00240.x>
- Lacker, J. M. & Weinberg, J. A. (1989). Optimal contracts under costly state falsification. *Journal of Political Economy* 97(6):1345–1363. <https://doi.org/10.1086/261657>
- Lazear, E. P. (2006). Speeding, terrorism, and teaching to the test. *Quarterly Journal of Economics* 121(3):1029–1061.
- Lin, C.-c. (2004). Bonding, shirking and adverse selection. *Economic Modelling* 21(3):545–560.
- Maggi, G. & Rodríguez-Clare, A. (1995). Costly distortion of information in agency problems. *RAND Journal of Economics* 26(4):675–689. <https://doi.org/10.2307/2556012>
- Manelli, A. M. & Vincent, D. R. (1995). Optimal procurement mechanisms. *Econometrica* 63(3):591–620.
- Mookherjee, D. & Png, I. (1989). Optimal auditing, insurance, and redistribution. *Quarterly Journal of Economics* 104(2):399–415. <https://doi.org/10.2307/2937855>
- Mookherjee, D. & Png, I. P. L. (1994). Marginal deterrence in enforcement of law. *Journal of Political Economy* 102(5):1039–1066.
- Mylovanyov, T. & Zapechelyuk, A. (2017). Optimal allocation with ex post verification and limited penalties. *American Economic Review* 107(9):2666–2694. <https://doi.org/10.1257/aer.20140494>

- Patterson, E. R. & Noel, J. (2003). Audit strategies and multiple fraud opportunities of misreporting and defalcation. *Contemporary Accounting Research* 20(3):519–549.
- Perez-Richet, E. & Skreta, V. (2022). Test design under falsification. *Econometrica* 90(3):1109–1142. <https://doi.org/10.3982/ECTA16346>
- Perez-Richet, E. & Skreta, V. (2024). Score-based mechanisms. Working paper. <https://arxiv.org/abs/2403.08031>
- Polinsky, A. M. & Shavell, S. (2000). The economic theory of public enforcement of law. *Journal of Economic Literature* 38(1):45–76. <https://doi.org/10.1257/jel.38.1.45>
- Qiu, X. & Shan, L. (2025). Multi-dimensional test design. Working paper. <https://arxiv.org/abs/2502.12264>
- Rhoades, S. C. (1999). The impact of multiple component reporting on tax compliance and audit strategies. *The Accounting Review* 74(1):63–85.
- Şabac, F. & Yoo, J. (2018). Performance measure aggregation in multi-task agencies. *Contemporary Accounting Research* 35(2):716–733. <https://doi.org/10.1111/1911-3846.12418>
- Shapiro, C. & Stiglitz, J. E. (1984). Equilibrium unemployment as a worker discipline device. *American Economic Review* 74(3):433–444. <https://www.jstor.org/stable/1804018>
- Shavell, S. (1992). A note on marginal deterrence. *International Review of Law and Economics* 12(3):345–355.
- Stiglitz, J. E. & Weiss, A. (1981). Credit rationing in markets with imperfect information. *American Economic Review* 71(3):393–410.
- Thompson, K. (1984). Reflections on trusting trust. *Communications of the ACM* 27(8):761–763.
- Townsend, R. M. (1979). Optimal contracts and competitive markets with costly state verification. *Journal of Economic Theory* 21(2):265–293. [https://doi.org/10.1016/0022-0531\(79\)90031-0](https://doi.org/10.1016/0022-0531(79)90031-0)
- Yang, F. (2025). Costly multidimensional screening. *Review of Economic Studies*, advance article rdaf105 (published online 24 December 2025). <https://doi.org/10.1093/restud/rdaf105>

Institutional and market facts cited in Section 6 (SecNumCloud v3.2 criteria and qualification record; the EUCS March-2024 draft; the Microsoft EU Data Boundary and the 10 June 2025 French Senate testimony; AWS European Sovereign Cloud) are sourced and dated in the project’s fact-check log; the last full verification against primary sources was completed on 3 July 2026.