

PREPRINT - NOT PEER REVIEWED

LAKANA SOS (Sovereign Operating System)

Attribution-Controlled Monte Carlo Evaluation of a Local-First Civilian Safety Architecture Under Synthetic Degraded-State Stress

Engineering Archive (engrXiv) submission edition

MarTaize K. Fails

Founder and Principal Architect, LAKANA Sovereign Systems LLC

Norman, Oklahoma, USA

Correspondence: martaize@lakana.systems

September 2026

Scope boundary

This paper reports simulation evidence. It does not establish field validation, deployment readiness, emergency-response certification, medical efficacy, measured attacker resistance, or empirically calibrated human-safety performance.

Numerical reporting policy: source-locked values are reported at their archived precision; rounded narrative substitutes are not used for load-bearing SOS outcomes.

Abstract

Civilian safety-architecture evaluations can confound detector behavior, delivery infrastructure, and full-stack policy composition, obscuring which subsystem produces a reported gain. This study evaluates LAKANA SOS, a local-first civilian safety architecture, through three separated Monte Carlo suites: shared-detector delivery/infrastructure, shared-transport detector trade-space, and end-to-end integration. The archived run records 250,000 nominal trial scenarios and 500,000 synthetic degraded-state trial scenarios. In the headline comparison, the simulator-defined protected-outcome proportion is 0.954636 for LAKANA + CivOS (95% Wilson interval [0.953813, 0.955445]), 0.949616 for LAKANA without CivOS ([0.948752, 0.950467]), and 0.947172 for a synthetic centralized/cloud comparator ([0.946288, 0.948042]). The LAKANA + CivOS minus comparator contrast is 0.007464, equivalent to 0.746400 percentage points, with a bias-corrected and accelerated 95% bootstrap interval [0.007128, 0.007800]. Under the nominal shared-detector suite, detector-side rates are identical across architecture arms; conditional delivery given escalation is 0.500460 for both LAKANA transport variants and 0.237438 for the comparator, while BFB delivery given need is 1.000000 for LAKANA + CivOS and 0.199586 for the comparator. Under synthetic degraded-state shared-detector stress, comparator protected outcome changes from 0.947000 to 0.782674 (absolute change -0.164326; relative change -17.352270%), whereas LAKANA + CivOS changes from 0.954842 to 0.950753 (absolute change -0.004089; relative change -0.428238%). The detector suite does not support detector dominance: more aggressive detector conditions achieve higher recall and protected outcome while imposing higher escalation and false-positive burden. Eight independent replications preserve the architecture ordering. Sobol total-order sensitivity for the CivOS-minus-comparator delta is concentrated on detector aggressiveness / TSARO threshold ($ST = 0.992734$) within the bounded six-parameter analysis region. The anytime-precision procedure reaches the 5,000,000-sample cap without satisfying its 0.001000 fail-closed target. The supported conclusion is therefore bounded: the simulation provides attribution-controlled evidence of delivery- and fail-closed architecture advantages under modeled conditions, not field proof.

Keywords: civilian safety architecture; local-first systems; fail-closed systems; Monte Carlo simulation; degraded communications; cyber-physical safety; evidence integrity; TSARO; NICOLE Protocol; CivOS; LAKANA SOS

Contribution statement

- Separates delivery/infrastructure effects from detector operating-point effects and from full-stack composition.
- Reports exact source-locked protected-outcome, confidence-interval, delivery, burden, replication, sensitivity, and numerical-diagnostic values.
- Shows that the strongest architecture result occurs in a shared-detector comparison in which detector-side behavior is held constant.
- Preserves the detector result as a burden-sensitivity trade-space instead of claiming universal detector superiority.
- Makes provenance defects, synthetic stress assumptions, failed anytime stopping, and non-field status explicit rather than treating them as hidden caveats.

1. Introduction

Civilian safety systems increasingly combine local sensors, communication links, remote services, evidence handling, and automated escalation. When those components are evaluated only as one integrated endpoint, a favorable outcome does not identify whether the gain came from sensing, trigger aggressiveness, transport redundancy, fail-closed execution, or downstream integrity logic. That attribution problem is especially consequential for safety systems intended to operate when power, radio-frequency conditions, network reachability, or external services are degraded.

LAKANA SOS means the LAKANA Sovereign Operating System for civilian safety. The architecture is designed around local-first state, bounded authority, fail-closed behavior, evidence custody, and multi-path communication. In the public-safe architecture, TSARO represents the threat-adaptive bounded risk and escalation layer; NICOLE Protocol represents evidence-integrity, separation, and auditability functions; CivOS represents the local survival substrate that arbitrates persistent state, energy, and protected execution; and Blue Force Bridge (BFB) represents a bounded responder-handoff pathway. The Monte Carlo model evaluates abstractions of these functions. It is not a field implementation of emergency dispatch, responder arrival, or physical rescue.

The central methodological decision in this paper is separation. A shared-detector suite fixes detector-side behavior while transport and CivOS-mediated infrastructure vary. A shared-transport suite fixes transport while detector operating point varies. An end-to-end suite permits full-stack composition to vary but is treated as attribution-limited. Synthetic degraded-state reruns then ask whether the same relationships persist when modeled RF hostility and compromise conditions are increased.

The numerical authority for this submission is the LAKANA SOS Public Release v1 separated-suite output lineage. Historical v71 materials were reviewed for terminology and claim-boundary context but are not used as the numerical authority for this manuscript. September 2026 public white-paper materials were reviewed for current SOS naming and governance language; they are not substituted for the archived Monte Carlo result artifacts.

1.1 Bounded research questions

1. RQ1: With detector behavior held constant, do delivery/infrastructure and CivOS-mediated pathways change the simulator-defined protected outcome?
2. RQ2: With transport held constant, how do detector operating points trade recall and protected outcome against escalation and false-positive burden?
3. RQ3: Under synthetic degraded-state stress, does the architecture ordering observed in nominal conditions persist?
4. RQ4: Are the headline ordering and deltas stable across independent seeds, bounded sensitivity analysis, and numerical diagnostics?

2. Architecture and evaluation boundary

Layer / object	Public-safe role	Representation in this study
LAKANA SOS	User-facing civilian safety operating architecture	Full architecture labels and end-to-end comparison suites
TSARO	Threat-adaptive bounded risk and escalation logic	Detector operating points, recall, false-positive rate, escalation rate, threshold sensitivity
NICOLE Protocol	Evidence integrity, separation, and auditability	Integrity scores, replay/stream compromise proxies, separation diagnostics
CivOS	Local survival substrate for protected state, energy, and fail-closed operation	Ablation comparison, protected fallback, energy and OS-compromise diagnostics
Blue Force Bridge (BFB)	Bounded responder-handoff pathway	BFB delivery given need

Table 1. Public-safe architecture mapping. The study evaluates modeled architecture functions, not deployment-certified hardware or emergency-service integration.

The exported label “survivability” is retained in artifact names but interpreted here as a simulator-defined protected-outcome proportion. It is not mortality, clinical survival, injury prevention, responder success, or field safety. This distinction governs every result below.

3. Study provenance, source authority, and computation

3.1 Frozen execution lineage

The public-release provenance record preserves a code-manifest path in the v47 lineage and a runtime-manifest internal label of v30. The mismatch is retained rather than rewritten. The public release notes identify it as a stale internal label while preserving the underlying execution names for auditability. Numerical outputs, trial allocations, figures, and result tables are treated as authoritative only where they reconcile across the archived result files.

Field	Source-locked value
Public package label	LAKANA SOS Public Release v1
Code-manifest lineage	lakana_sos_apex_mc_v47.py
Code SHA-256 prefix	5b7266eb236c48a3...
Runtime internal label	v30
Nominal trial scenarios	250,000
Synthetic degraded-state trial scenarios	500,000
Workers	60
Batch size	25,000
Recorded runtime	844.57 s
empirical_enable	false
env_trace_enable	false
Python	3.11.2
NumPy / SciPy / pandas / matplotlib	2.4.3 / 1.17.1 / 3.0.1 / 3.10.8
Platform	Linux 6.1.0-44-cloud-amd64 / glibc 2.36

Table 2. Focused provenance record for the current numerical authority.

No empirical calibration or external environmental trace ingestion is claimed for this run. The degraded-state regime is synthetic. The broader LAKANA Monte Carlo research program has used Amazon Web Services and Google Cloud computational infrastructure; the public runtime record for this specific archived execution identifies a Linux cloud environment but does not encode a provider identity, so no provider is attributed to this exact run in the results section.

3.2 Artifact authority order

Claim family	Controlling artifact family
Headline architecture outcome	survivability_summary_all_arch.json
Cross-architecture inference	bootstrap_and_tost.json
Nominal suites	publication_delivery_suite.json; publication_detection_suite.json; publication_end_to_end_suite.json
Synthetic degraded-state suites	publication_delivery_suite_adv.json; publication_detection_suite_adv.json; publication_end_to_end_suite_adv.json
Failure-mode audit	failure_code1_diagnostics.json
Replication	replication_survivability.csv; replication_robustness_summary.json
Sensitivity	Sobol sensitivity outputs for CivOS-industry and CivOS-no-CivOS deltas
Numerical diagnostics	spd_projection_diagnostics.json; anytime_precision_stopping.json

Table 3. Artifact authority used to reconcile manuscript prose and tables.

3.3 Version reconciliation performed for this submission

The engrXiv submission edition does not merge incompatible Monte Carlo branches. The Public Release v1 separated-suite outputs control all numerical claims. Historical v71 outcomes are excluded from quantitative synthesis. A stale adversarial end-to-end appendix row present in one manuscript rendering was replaced with the controlling output values: 0.783614, 0.831050, 0.952107, and 0.980671. Those values match both the detailed supplement and the source publication-suite table. No numerical averaging across versions was performed.

4. Methods

4.1 Separated suite design

Suite	Controlled factor	Varying factor	Nominal n	Stress n
Shared-detector delivery	Detector behavior	Transport, infrastructure, CivOS	95,000	190,000
Shared-transport detector	Transport	Detector operating point	85,000	170,000
End-to-end integrated	None	Full-stack composition	70,000	140,000

Table 4. Suite design and exact trial allocations.

The suite structure constrains causal language. Shared-detector comparisons can support delivery/infrastructure interpretations because the detector-side rates are fixed by construction. Shared-transport comparisons can support detector operating-point interpretations because transport is fixed. End-to-end comparisons are operational summaries and do not independently identify subsystem causality.

4.2 Outcome definitions

Metric	Interpretation
Protected outcome / “survivability”	Simulator-defined binary protected-outcome proportion; not clinical or field survival.
Recall given threat	Detector-side conditional success rate in threat-present trials.
False-positive rate	Escalation rate in non-threat trials.
Precision given escalation	Positive predictive value of escalations under the simulation definition.
Conditional delivery escalation	Delivery success conditioned on an escalation; delivery/infrastructure metric.
BFB delivery need	Modeled BFB delivery fraction among simulator-tagged need cases.
Failure-code-1 rate	Principal exported miss/audit class used for diagnostic decomposition.
Integrity mean	Simulator-internal integrity quantity; not independent cryptographic certification.

Table 5. Exported metric interpretation used throughout this submission.

4.3 Statistical reporting

For architecture a , the headline estimator is the sample protected-outcome proportion $S_a = n^{-1} \sum_i I(\text{outcome}_i = \text{protected})$. Per-arm uncertainty is reported with 95% Wilson intervals. Cross-architecture inference uses the archived nonparametric bootstrap contrast $\Delta(a,b) = S_a - S_b$. The primary contrast uses $B = 10,000$ bootstrap resamples. The BCa correction for LAKANA + CivOS minus the synthetic centralized/cloud comparator uses $z_0 = -0.0248$ and $a = 0.0038$ in the archived result.

The archived TOST calculation uses a simulator-scale equivalence margin of 0.010000 and standard error 0.000172. It is retained only as a secondary diagnostic because the selected equivalence margin is not a field-calibrated safety margin. The load-bearing cross-architecture interval is the BCa bootstrap interval.

4.4 Dependence model and numerical repair

The transport dependence model uses five modeled channels - BLE, UWB, WiFi NAN, LoRa, and Satellite - with an equicorrelation coefficient $\rho = 0.550000$. The sensor dependence matrix is heterogeneous, with off-diagonal entries from 0.050000 to 0.450000. Both are projected to the nearest symmetric positive semi-definite matrix before copula sampling. Frobenius-relative distortions are $1.06e-16$ for transport and $3.47e-16$ for sensors. These values support numerical feasibility of the repair step; they do not validate the assumed dependence structure against field measurements.

4.5 Robustness diagnostics

Robustness is assessed by eight independent 20,000-trial replications, a threat-stratified consistency check, bounded Sobol sensitivity, and an anytime-precision diagnostic. The Sobol design uses $N = 2,048$ base samples, six parameter families, 3,048 inner Monte Carlo trials per design point, scrambling enabled, and convergence prefixes at $N = 1,024$ and $N = 2,048$. The anytime procedure is fail-closed: its configured target is $\epsilon = 0.001000$ and it is not declared successful unless both Hoeffding and empirical-Bernstein half-width conditions are satisfied.

5. Results

5.1 Headline architecture comparison

Architecture	Protected outcome	Wilson low	Wilson high
Industry centralized/cloud comparator	0.947172	0.946288	0.948042
LAKANA without CivOS	0.949616	0.948752	0.950467
LAKANA + CivOS	0.954636	0.953813	0.955445

Table 6. Headline architecture-level protected-outcome summary (archived $n = 250,000$).

LAKANA + CivOS exceeds the synthetic centralized/cloud comparator by 0.007464, or 0.746400 percentage points. The 95% BCa interval is $[0.007128, 0.007800]$, and the percentile interval is $[0.007132, 0.007808]$. LAKANA + CivOS exceeds LAKANA without CivOS by 0.005020, or 0.502000 percentage points. These are simulator-defined contrasts, not field effect sizes.

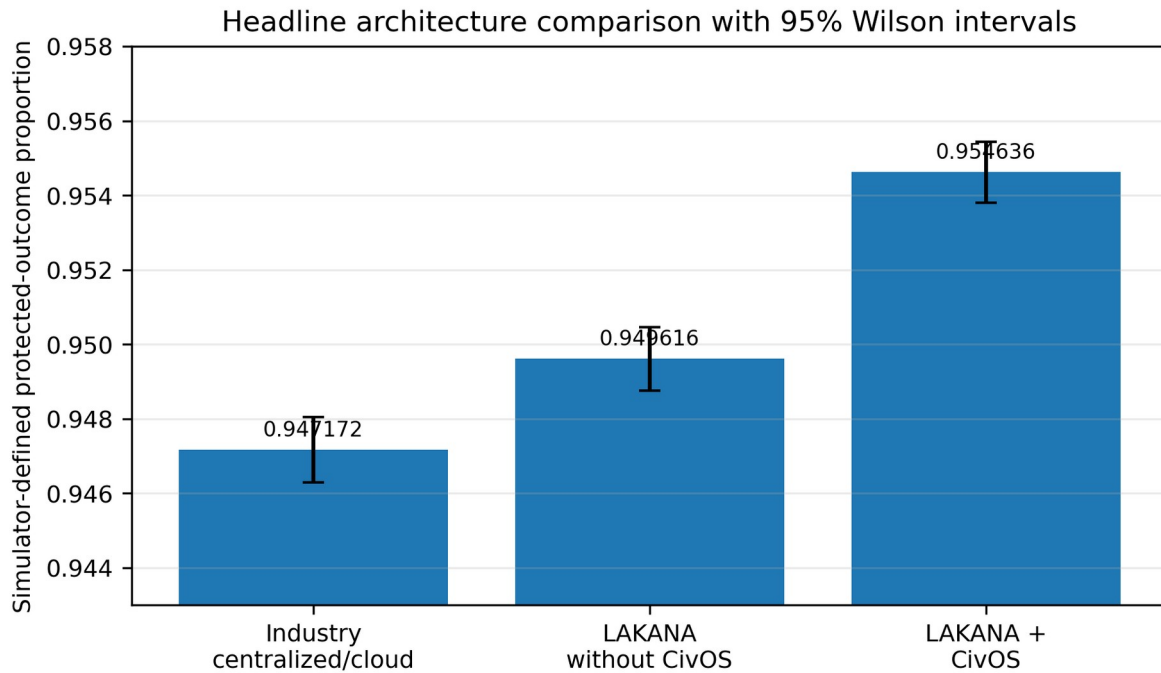


Figure 1. Headline architecture comparison. Data labels preserve six-decimal source precision; error bars are 95% Wilson intervals.

5.2 Nominal shared-detector delivery/infrastructure suite

Architecture	Protected outcome	Cond. delivery esc.	Cond. protected esc.	BFB need
Industry centralized/cloud comparator	0.947000	0.237438	0.862875	0.199586
LAKANA without CivOS	0.949653	0.500460	0.909258	0.434392
LAKANA + CivOS	0.954842	0.500460	1.000000	1.000000

Table 7. Nominal shared-detector suite ($n = 95,000$).

Recall (0.179575), false-positive rate (0.050061), escalation rate (0.057189), precision given escalation (0.172833), failure-code-1 rate (0.045158), and integrity mean (0.890745) are identical across all three architecture arms by construction. Conditional delivery given escalation is 0.500460 for both LAKANA transport variants and 0.237438 for the comparator, a ratio of 2.107750. The supported attribution is downstream delivery/infrastructure and protected-outcome conversion, not superior detector-side behavior.

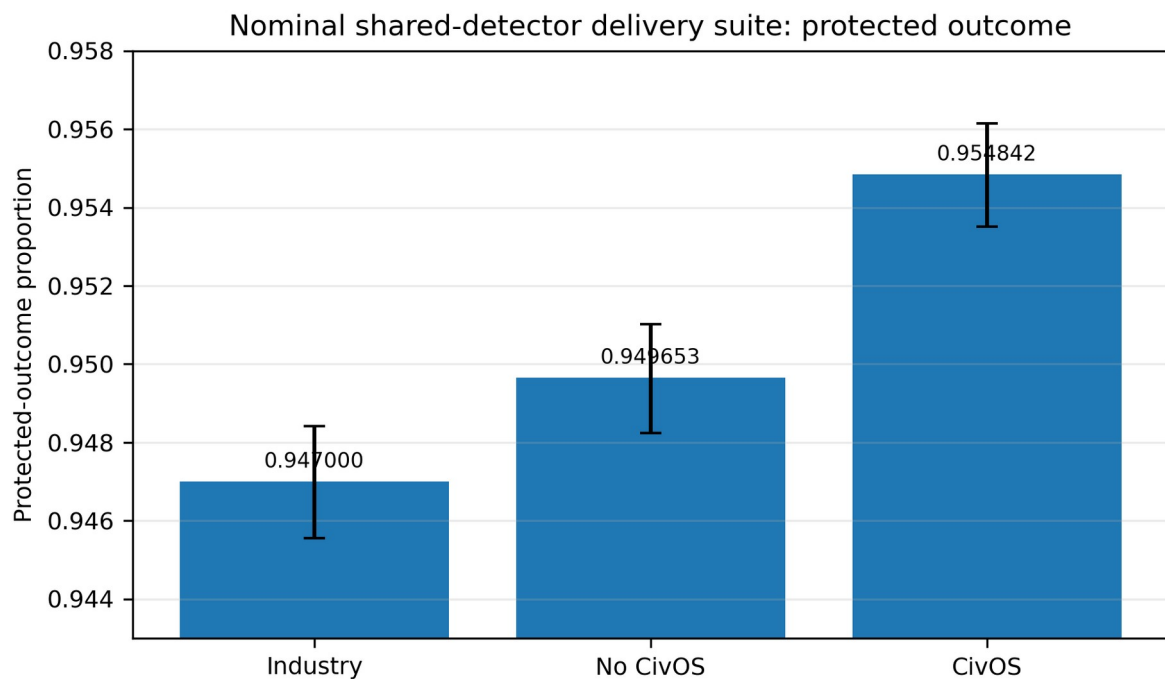


Figure 2. Nominal shared-detector protected outcome with 95% Wilson intervals.

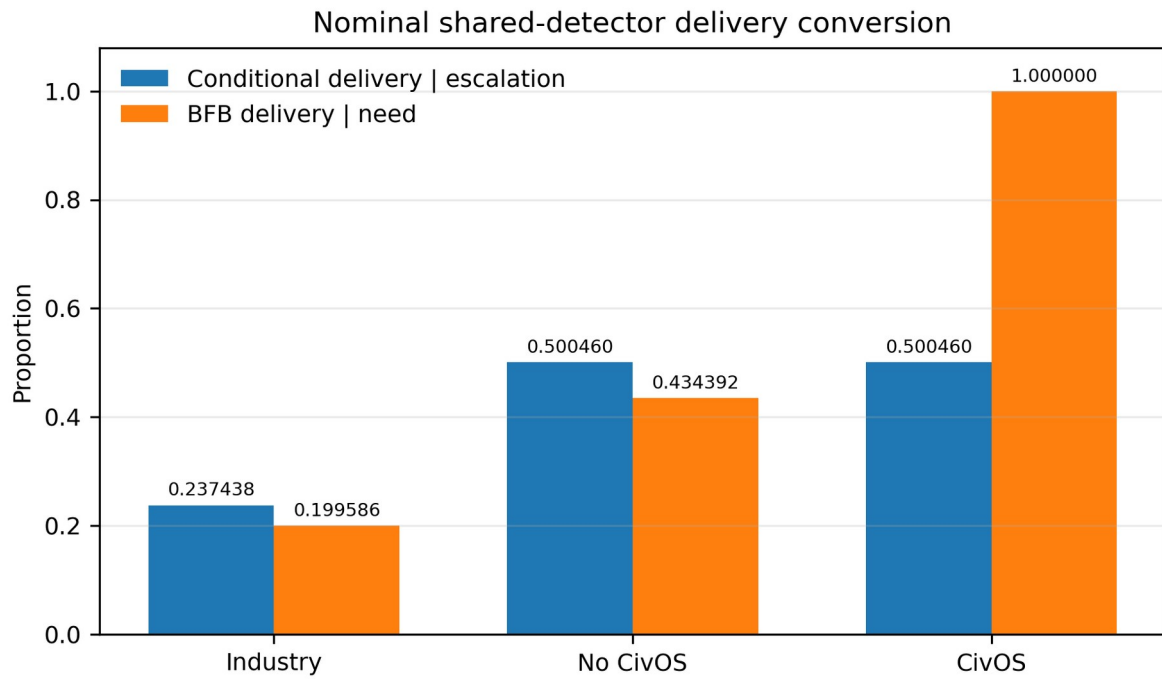


Figure 3. Nominal shared-detector conditional delivery and BFB delivery. Values are source-locked to six decimals.

5.3 Nominal shared-transport detector trade-space

Condition	Protected outcome	Recall threat	FP rate	Escalation rate
Detector floor	0.946353	0.000000	0.000000	0.000000
Full TSARO	0.955471	0.169956	0.052797	0.059082
No-context detector	0.968071	0.404825	0.093784	0.110471
Naive detector	0.984588	0.713377	0.094804	0.127988

Table 8. Nominal shared-transport detector suite ($n = 85,000$).

The naive detector reaches protected outcome 0.984588 and recall 0.713377, compared with 0.955471 and 0.169956 for Full TSARO. The naive operating point also raises escalation from 0.059082 to 0.127988, a ratio of 2.166277, and raises the false-positive rate from 0.052797 to 0.094804, a ratio of 1.795632. The result therefore defines a burden-sensitivity trade-space; it does not establish detector dominance for Full TSARO.

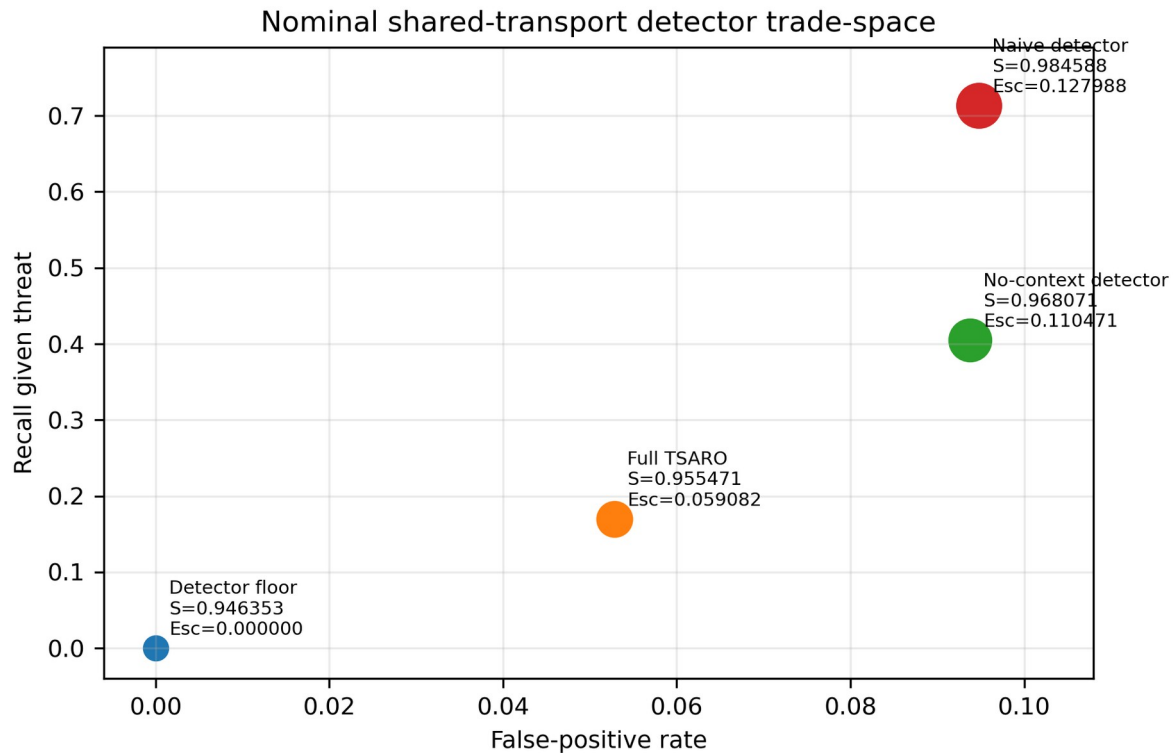


Figure 4. Nominal detector trade-space under shared transport. Marker area scales with escalation rate; annotations report exact protected outcome and escalation rate.

5.4 Nominal end-to-end integrated suite

Condition	Protected outcome	Wilson low	Wilson high	Recall	FP rate	BFB need
Industry shared-detector condition	0.945814	0.944113	0.947467	0.180407	0.050295	0.162059
LAKANA without CivOS	0.948471	0.946809	0.950085	0.180407	0.050295	0.409914
LAKANA full CivOS	0.953986	0.952409	0.955513	0.180407	0.050295	1.000000
LAKANA naive detector + CivOS	0.983557	0.982588	0.984473	0.707125	0.093522	1.000000

Table 9. Nominal end-to-end integrated suite ($n = 70,000$).

Because detector and infrastructure vary simultaneously in this suite, the end-to-end values are operational summaries. They do not overturn the controlled-suite interpretation: the naive detector + CivOS condition has the highest protected outcome, but it also inherits the higher detector burden documented in Table 8.

5.5 Synthetic degraded-state shared-detector results

Architecture	Protected outcome	Wilson low	Wilson high	Cond. delivery esc.	BFB need
Industry centralized/cloud comparator	0.782674	0.780813	0.784522	0.184382	0.136057
LAKANA without CivOS	0.828958	0.827258	0.830644	0.434117	0.373596
LAKANA + CivOS	0.950753	0.949771	0.951717	0.434117	1.000000

Table 10. Synthetic degraded-state shared-detector suite ($n = 190,000$).

The synthetic centralized/cloud comparator changes from 0.947000 nominal protected outcome to 0.782674, an absolute change of -0.164326 and a relative change of -17.352270%. LAKANA + CivOS changes from 0.954842 to 0.950753, an absolute change of -0.004089 and a relative change of -0.428238%. The shared-detector CivOS-minus-comparator separation therefore changes from 0.007842 nominally to 0.168079 under the synthetic degraded-state regime.

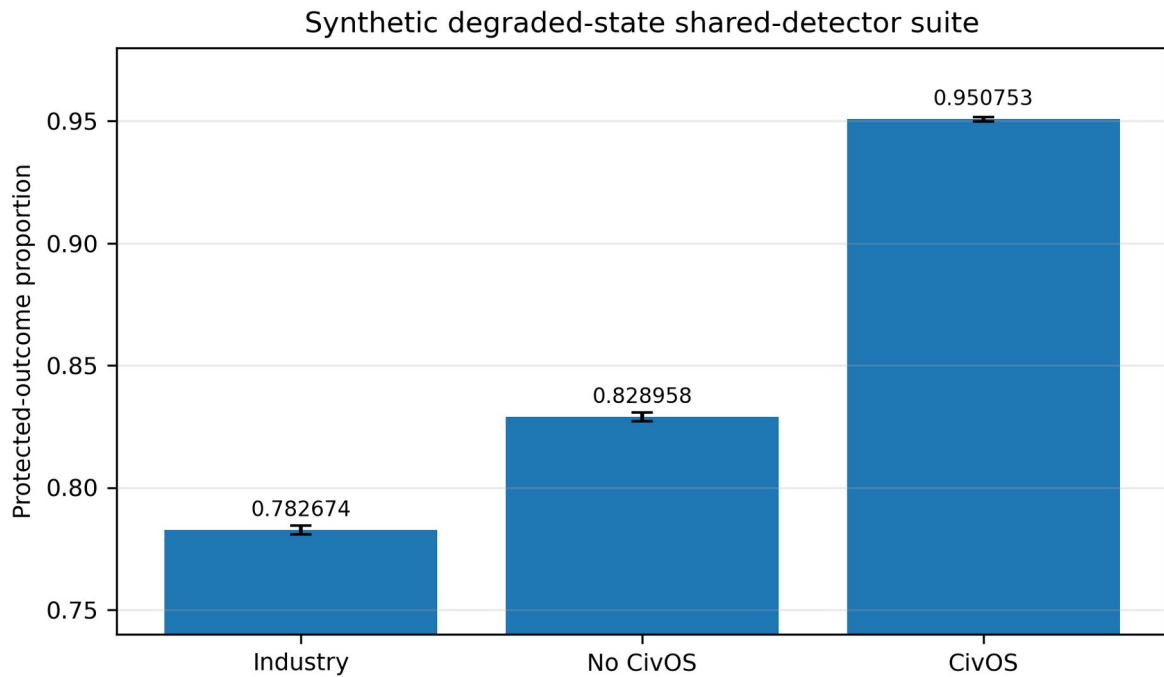


Figure 5. Synthetic degraded-state shared-detector protected outcome with 95% Wilson intervals.

5.6 Synthetic degraded-state detector trade-space

Condition	Protected outcome	Recall threat	FP rate	Escalation rate
Detector floor	0.944688	0.000000	0.000000	0.000000
Full TSARO	0.950818	0.111028	0.029484	0.033994
No-context detector	0.963553	0.341806	0.070431	0.085441
Naive detector	0.980753	0.653302	0.067629	0.100024

Table 11. Synthetic degraded-state shared-transport detector suite ($n = 170,000$).

The detector burden-sensitivity pattern persists under stress. Full TSARO reaches protected outcome 0.950818 with recall 0.111028, false-positive rate 0.029484, and escalation rate 0.033994. The naive detector reaches protected outcome 0.980753 with recall 0.653302, false-positive rate 0.067629, and escalation rate 0.100024. These values reinforce the trade-space interpretation rather than a universal detector-superiority claim.

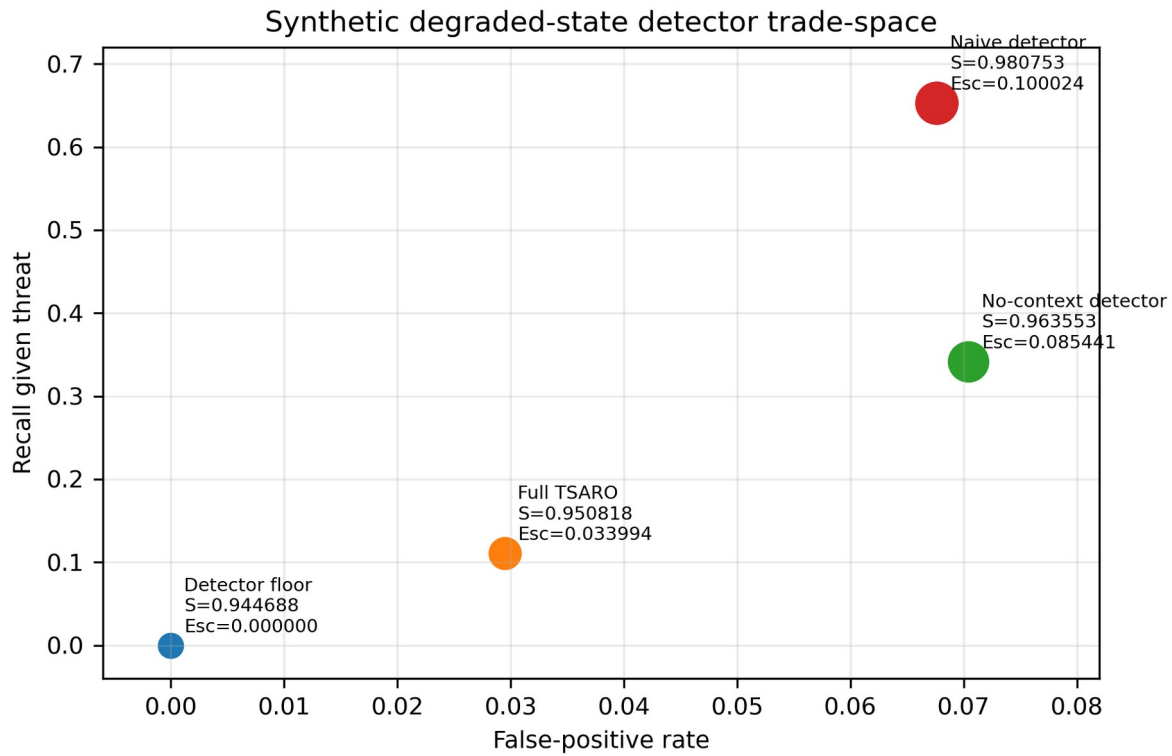


Figure 6. Synthetic degraded-state detector trade-space. Marker area scales with escalation rate; annotations preserve six-decimal source precision.

5.7 Synthetic degraded-state end-to-end integrated suite

Condition	Protected outcome	Wilson low	Wilson high	Recall	FP rate	BFB need
Industry shared-detector condition	0.783614	0.781449	0.785764	0.122743	0.030486	0.132130
LAKANA without CivOS	0.831050	0.829078	0.833004	0.122743	0.030486	0.376282
LAKANA full CivOS	0.952107	0.950976	0.953213	0.122743	0.030486	0.999966
LAKANA naive detector + CivOS	0.980671	0.979937	0.981380	0.647344	0.069327	0.999966

Table 12. Synthetic degraded-state end-to-end integrated suite ($n = 140,000$). These values are the controlling output-table values used to correct a stale appendix row in an earlier manuscript rendering.

The full CivOS end-to-end condition reaches 0.952107 while the synthetic centralized/cloud condition reaches 0.783614. The naive detector + CivOS condition reaches 0.980671 with recall 0.647344 and false-positive rate 0.069327. Because this suite varies both detector and infrastructure, these values remain attribution-limited.

5.8 Failure-mode transparency and integrity diagnostics

Failure-code-1 occurs in 11,341 of the 250,000 nominal headline trial scenarios, for an exported rate of 0.045364. The diagnostic decomposition is reported because it makes the simulator's degraded outcomes inspectable. It is not evidence that the labeled subclasses are empirically validated physical causes.

Subclass	Count	Rate given code 1
FDE suppression	5,499	0.484878
Context interference	1,576	0.138965
Bounded suppression	1,433	0.126356
Mid threshold	1,064	0.093819
Far below threshold	1,032	0.090997
Near threshold	580	0.051142
Confirmation suppressed	135	0.011904
Integrity deadzone	17	0.001499
Hard veto	5	0.000441

Table 13. Exact failure-code-1 subclass decomposition.

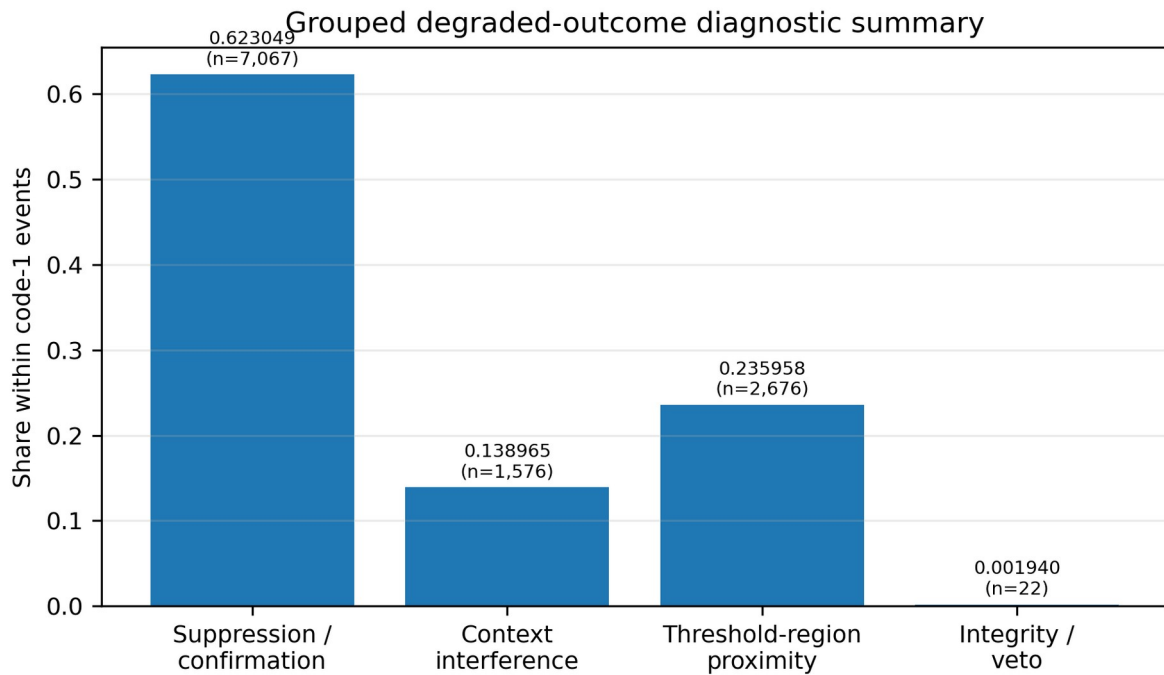


Figure 7. Public-safe grouped failure-code-1 diagnostic summary. Group shares and counts are exact aggregations of the nine source subclasses.

Integrity / fail-closed metric	Value
Evidence integrity rate	0.999884
Replay compromise proxy	0.000040
Stream compromise proxy	0.000076
Maximum key-compromise proxy	0.002580
Battery initial mean	0.549126
Battery after core mean	0.548297
Core energy cost mean	0.000829
Protected fallback success rate	0.999996
Fallback survival given OS compromise	0.995000

Table 14. Bounded simulator-internal integrity and fail-closed summary.

The integrity table supports internal consistency claims only. It does not establish hardware assurance, independent cryptographic certification, regulatory sufficiency, or resilience against real attackers.

6. Robustness, sensitivity, and numerical diagnostics

6.1 Eight-seed replication

Rep.	Seed	Industry	No CivOS	CivOS
0	4222882316	0.947650	0.950000	0.955000
1	2563681557	0.945300	0.947800	0.953000
2	3614085311	0.948300	0.950800	0.955850
3	1330953553	0.944350	0.946850	0.952750
4	3916710891	0.946350	0.948850	0.953900
5	2606998623	0.948550	0.951150	0.956500
6	2418506523	0.948400	0.950100	0.956500
7	609688932	0.945700	0.948050	0.953500

Table 15. Replication-by-seed protected-outcome proportions (n = 20,000 per replication).

Architecture	Mean	SD	MCSE of mean
Industry centralized/cloud comparator	0.946825	0.001615	0.000571
LAKANA without CivOS	0.949200	0.001546	0.000547
LAKANA + CivOS	0.954625	0.001541	0.000545

Table 16. Exact eight-seed replication summary.

All eight replications preserve the same ordering: LAKANA + CivOS > LAKANA without CivOS > synthetic centralized/cloud comparator. This supports seed-level stability for the configured model; it does not establish robustness to arbitrary structural assumptions.

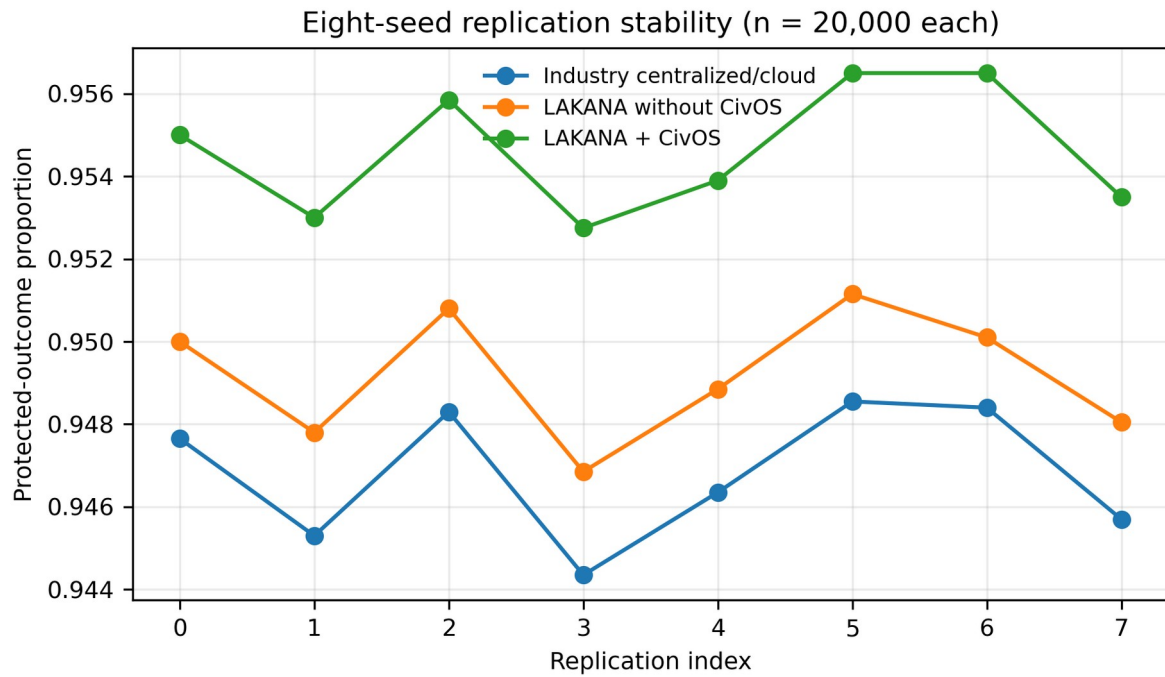


Figure 8. Protected-outcome stability across eight independent seeds.

6.2 Threat-stratified consistency check

Architecture	Stratified estimate	Difference from headline
Industry centralized/cloud comparator	0.947217	+0.000045
LAKANA without CivOS	0.949774	+0.000158
LAKANA + CivOS	0.954848	+0.000212

Table 17. Threat-stratified diagnostic estimates and exact difference from the headline estimator.

The stratified estimator remains a diagnostic. The fixed-N headline architecture outputs remain the primary estimators.

6.3 Sobol sensitivity

Parameter family	S1	ST
Adversarial effort	0.000052	0.002215
OS compromise probability	0.000365	0.001256
HW compromise probability	-0.000367	0.000031
Initial energy reserve	0.000000	0.000000
Detector aggressiveness / TSARO threshold	0.978753	0.992734
Transport degradation	0.006315	0.019217

Table 18. Sobol indices for Delta(LAKANA + CivOS - synthetic centralized/cloud comparator).

Parameter family	S1	ST
Adversarial effort	0.002595	0.012067
OS compromise probability	-0.000274	0.000841
HW compromise probability	0.000048	0.000019
Initial energy reserve	0.000000	0.000000
Detector aggressiveness / TSARO threshold	0.857517	0.952468
Transport degradation	0.046856	0.133457

Table 19. Sobol indices for Delta(LAKANA + CivOS - LAKANA without CivOS).

Within the bounded six-parameter analysis region, detector aggressiveness / TSARO threshold dominates both deltas, with ST = 0.992734 for CivOS minus comparator and ST = 0.952468 for CivOS minus no-CivOS. Transport degradation is secondary, with ST = 0.019217 and ST = 0.133457 respectively. Small negative S1 estimates are interpreted as numerical estimation noise around zero, not negative physical influence.

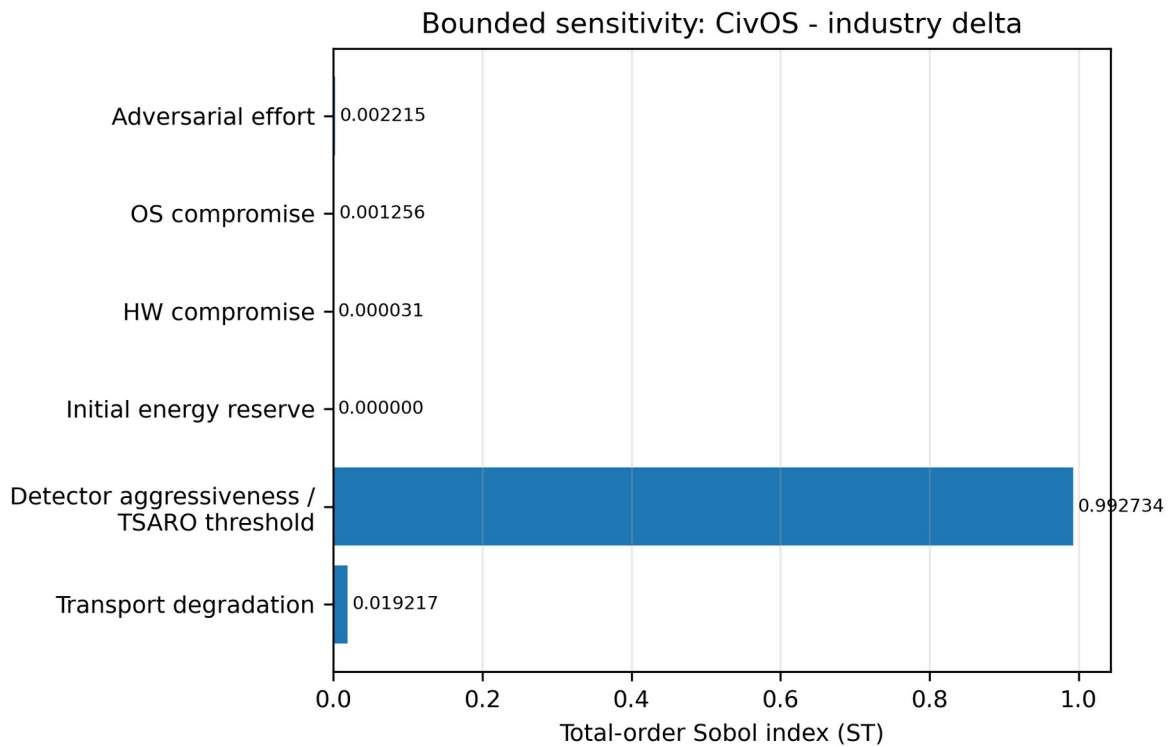


Figure 9. Total-order Sobol sensitivity for the CivOS-minus-comparator delta within the configured six-parameter box.

6.4 Anytime-precision stopping diagnostic

Metric	Mean	Hoeffding HW	Emp.-Bern. HW	Target
Delta: CivOS - industry	0.007683	0.002403	0.001215	0.001000
Delta: CivOS - no CivOS	0.005132	0.002403	0.001215	0.001000
Industry centralized/cloud comparator	0.946937	0.001202	0.000545	0.001000
LAKANA + CivOS	0.954620	0.001202	0.000507	0.001000
LAKANA without CivOS	0.949488	0.001202	0.000533	0.001000

Table 20. Anytime-precision status at the 5,000,000-sample cap.

The stopping criterion is not met. The per-arm empirical-Bernstein half-widths fall below 0.001000, but the corresponding Hoeffding half-widths remain 0.001202, and both delta metrics remain above the target under both bound families. The procedure is therefore a failed stopping attempt and a transparency diagnostic, not a convergence certificate.

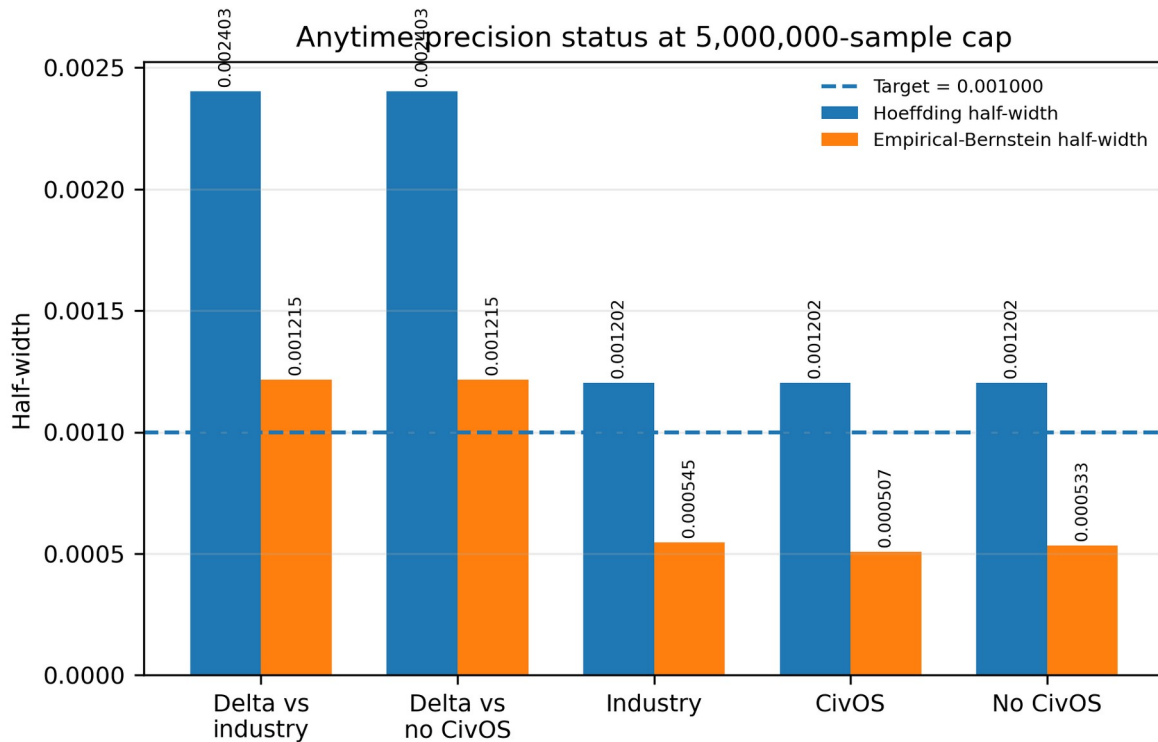


Figure 10. Final confidence-sequence half-widths at the 5,000,000-sample cap. The dashed target is 0.001000.

6.5 SPD projection diagnostic

The nearest-SPD projection changes the transport and sensor correlation matrices by Frobenius-relative distortions of $1.06\text{e-}16$ and $3.47\text{e-}16$ respectively. This shows that numerical repair is negligible at floating-point scale. It does not validate the physical realism of the assumed correlation matrices.

7. Discussion

7.1 Strongest supported result

The cleanest result is architecture-level and delivery-mediated. In the nominal shared-detector suite, all detector-side rates are identical while protected outcome, conditional delivery, conditional protected outcome, and BFB delivery differ. Under synthetic degraded-state stress, the same controlled comparison widens from a CivOS-minus-

comparator protected-outcome gap of 0.007842 to 0.168079. Because detector behavior is fixed in this suite, the result is consistent with a delivery/infrastructure and fail-closed contribution inside the simulation.

7.2 Detector result is a trade-space

The detector results prevent a stronger but unsupported narrative. More aggressive detector conditions achieve materially higher recall and protected outcome while increasing false positives and escalation. Full TSARO is therefore characterized as a more selective operating point, not as a universally superior detector. Any deployment-oriented detector choice would require explicit burden constraints, human-factors evaluation, and field calibration that are absent from the current run.

7.3 What the stress regime establishes and does not establish

The synthetic degraded-state regime is useful as an internal stress test because it tests ordering under a harder configured environment. It is not an empirically calibrated attacker model, does not estimate real adversary economics, and does not prove that the same degradation magnitudes would occur in field networks. The correct inference is conditional: under this modeled stress regime, the centralized/cloud comparator loses substantially more protected outcome than the full CivOS variant.

7.4 Sensitivity constrains interpretation

The Sobol results show that detector aggressiveness / TSARO threshold dominates variance in the main architecture deltas inside the chosen six-parameter box. That finding strengthens the need for the separated-suite design: an integrated result alone could be misleading if threshold choice drives most output variance. The sensitivity result is local to the configured parameter region and cannot be promoted to a universal real-world importance ranking.

8. Limitations and non-claims

- No field validation. The study contains no live responder integration, human-safety outcome calibration, production hardware validation, or emergency-service certification.
- Synthetic degraded-state model. The stress regime is a designed simulation condition rather than an observed attacker or outage distribution.
- Synthetic comparator. The centralized/cloud comparator is a modeled architecture class and is not a measurement of every commercial safety platform.
- Assumed dependence structure. Transport and sensor correlation matrices are physics-prior assumptions, not field-measured channel correlations.
- Detector trade-space. Full TSARO is not the top performer on raw recall or protected outcome in the shared-transport detector suite.
- Anytime target not met. The 5,000,000-sample fail-closed precision target remains unsatisfied.
- Provenance label mismatch. The code-manifest v47 lineage and runtime internal v30 label are retained as an administrative imperfection.
- Controlled reproducibility. Public outputs are sufficient to inspect the reported aggregates but are intentionally insufficient to reconstruct protected implementation logic.
- No independent third-party code review is claimed for the protected simulation implementation.
- No medical, legal, law-enforcement, or guaranteed-rescue claim follows from the simulator-defined protected outcome.

9. Reproducibility and public-review posture

The submission is designed as a controlled-release engineering preprint. Public review can inspect exact aggregate outputs, Wilson intervals, bootstrap contrasts, suite allocations, replication values, failure-code decomposition, bounded Sobol indices, numerical-stability diagnostics, and the disclosed computation environment. The protected engine, exact internal control schedules, full implementation manifest, private chain-of-custody internals, and misuse-sensitive details are not reproduced in this paper.

The current submission package contains a source-locked values file used to regenerate the tables and figure annotations in this engrXiv edition. That file is an editorial audit aid and does not replace the archived execution artifacts. Any future public artifact deposit should preserve the original execution-lineage names rather than rewriting them to match a later publication label.

10. Conclusion

This paper presents an attribution-controlled Monte Carlo evaluation of LAKANA SOS under nominal and synthetic degraded-state conditions. The headline architecture comparison reports protected-outcome proportions of 0.954636 for LAKANA + CivOS, 0.949616 for LAKANA without CivOS, and 0.947172 for the synthetic centralized/cloud comparator, with a CivOS-minus-comparator contrast of 0.007464 and BCa 95% interval [0.007128, 0.007800]. The shared-detector suite provides the strongest architecture attribution because detector-side behavior is held constant while delivery and fail-closed outcomes differ. Under synthetic degraded-state stress, the shared-detector CivOS-minus-comparator gap is 0.168079. The detector suite simultaneously shows that aggressive detector settings buy higher recall and protected outcome with higher false-positive and escalation burden.

The evidence is therefore meaningful but bounded. It supports continued engineering development, external technical review, bench validation, and controlled pilot design. It does not establish field safety, deployment readiness, emergency-response performance, or real-world attacker resistance. The next evidentiary step is external calibration and hardware/field validation, not stronger language applied to the same simulation.

Statements and declarations

Author contribution

MarTaize K. Fails developed the LAKANA SOS architecture, simulation program, research questions, analytical framework, evidence boundaries, artifact-selection rules, interpretation, and publication decisions; performed or directed the computational work; reconciled the final quantitative claims against the archived output artifacts; and accepts responsibility for the manuscript.

Competing interests

MarTaize K. Fails is the founder and principal architect of LAKANA Sovereign Systems LLC and has direct intellectual-property and commercialization interests in LAKANA SOS, CivOS, TSARO, NICOLE Protocol, and related technologies. This relationship is disclosed as a competing financial and professional interest.

Funding and computational resources

No external research grant or sponsored-research support is claimed for this manuscript. The broader LAKANA Monte Carlo research program used Amazon Web Services and Google Cloud computing infrastructure. The public provenance record for the specific execution reported here identifies a Linux cloud environment but does not encode a provider identity; accordingly, neither provider is attributed to this exact run. The cloud providers had no claimed role in study design, interpretation, or publication decisions.

Ethics statement

This is a computational simulation study. No human participants, patient records, identifiable personal data, or live emergency events were used in the reported run. Human-subjects approval is therefore not claimed or represented as applicable to this simulation.

Data and code availability

Publication-facing aggregated results, manuscript tables, selected figures, and bounded reproducibility materials are available in the author-controlled public research package associated with this work. Full implementation source, protected thresholds, anti-abuse logic, private chain-of-custody internals, and reconstructive engine details are

withheld for intellectual-property, safety, and misuse-prevention reasons. Controlled technical review may be considered separately under appropriate confidentiality and responsible-use terms.

AI-assisted research-writing disclosure

OpenAI tools were used for editorial restructuring, drafting assistance, architecture presentation and polish, claims/non-claims consistency checking, formatting support, and adversarial editorial review. The author developed the architecture and simulation program, selected the research questions and evidence boundaries, controlled the source artifacts, verified the quantitative values against the archived outputs, reviewed the manuscript, and remains responsible for all claims, citations, analyses, limitations, and submission decisions. No AI system is listed as an author, and no AI-generated numerical result is treated as simulation evidence.

References

- [1] Demarta, S., & McNeil, A. J. (2005). The t copula and related copulas. *International Statistical Review*, 73(1), 111-129. <https://doi.org/10.1111/j.1751-5823.2005.tb00254.x>
- [2] Efron, B., & Tibshirani, R. J. (1993). *An Introduction to the Bootstrap*. Chapman and Hall/CRC. ISBN 9780412042317.
- [3] Higham, N. J. (1988). Computing a nearest symmetric positive semidefinite matrix. *Linear Algebra and its Applications*, 103, 103-118. [https://doi.org/10.1016/0024-3795\(88\)90223-6](https://doi.org/10.1016/0024-3795(88)90223-6)
- [4] Hoeffding, W. (1963). Probability inequalities for sums of bounded random variables. *Journal of the American Statistical Association*, 58(301), 13-30. <https://doi.org/10.1080/01621459.1963.10500830>
- [5] Howard, S. R., Ramdas, A., McAuliffe, J., & Sekhon, J. (2018). Time-uniform Chernoff bounds via nonnegative supermartingales. *arXiv:1808.03204*.
- [6] Jansen, M. J. W. (1999). Analysis of variance for model output. *Computer Physics Communications*, 117(1), 35-43. [https://doi.org/10.1016/S0010-4655\(98\)00154-4](https://doi.org/10.1016/S0010-4655(98)00154-4)
- [7] Maurer, A., & Pontil, M. (2009). Empirical Bernstein bounds and sample variance penalization. *arXiv:0907.3740*.
- [8] Saltelli, A., Annoni, P., Azzini, I., Campolongo, F., Ratto, M., & Tarantola, S. (2010). Variance based sensitivity analysis of model output: Design and estimator for the total sensitivity index. *Computer Physics Communications*, 181(2), 259-270. <https://doi.org/10.1016/j.cpc.2009.09.018>
- [9] Schuirmann, D. J. (1987). A comparison of the two one-sided tests procedure and the power approach for assessing the equivalence of average bioavailability. *Journal of Pharmacokinetics and Biopharmaceutics*, 15(6), 657-680. <https://doi.org/10.1007/BF01068419>
- [10] Sklar, A. (1959). Fonctions de repartition a n dimensions et leurs marges. *Publications de l'Institut de Statistique de l'Universite de Paris*, 8, 229-231.
- [11] Welford, B. P. (1962). Note on a method for calculating corrected sums of squares and products. *Technometrics*, 4(3), 419-420. <https://doi.org/10.1080/00401706.1962.10490022>
- [12] Wilson, E. B. (1927). Probable inference, the law of succession, and statistical inference. *Journal of the American Statistical Association*, 22(158), 209-212. <https://doi.org/10.1080/01621459.1927.10502953>
- [13] National Institute of Standards and Technology. (2022). *Engineering Trustworthy Secure Systems*. NIST SP 800-160 Vol. 1 Rev. 1. <https://doi.org/10.6028/NIST.SP.800-160v1r1>

Appendix A. Exact-value audit tables

A.1 Full nominal shared-detector detector-side equality

Metric	Industry	No CivOS	CivOS
Recall given threat	0.179575	0.179575	0.179575
False-positive rate	0.050061	0.050061	0.050061
Escalation rate	0.057189	0.057189	0.057189
Precision given escalation	0.172833	0.172833	0.172833
Failure-code-1 rate	0.045158	0.045158	0.045158
Integrity mean	0.890745	0.890745	0.890745

Table A1. Exact detector-side equality that licenses architecture attribution in the nominal shared-detector suite.

A.2 Exact integrity summary and computation endpoints

Exact reporting is deliberately retained in this submission. Values displayed as 1.000000 are source outputs at the available publication precision, not field guarantees. Values such as 0.000040 and 0.000076 are simulator-internal compromise proxies, not measured attack frequencies.