

AI ASSESS TECH

The Audit That Doesn't Wait for Quarter-End

Every audit you trust looks backward. Your AI is making consequential decisions right now. That blind spot is where AI risk lives — and where we work.

The blind spot every audit shares

A financial audit examines transactions that already happened. A SOC 2 report describes controls over a past period. An ISO certificate reflects a point-in-time review. Every audit discipline the enterprise trusts shares the same structure — and the same limitation: the period closes, then someone examines it. A deployed AI system doesn't work in periods. LLM agents make thousands of decisions between any two checkpoints, and its behavior drifts between them. The audit concept everyone respects has a structural blind spot exactly where AI risk lives. What would a year of AI drift costs you in terms of risk?

Same audit discipline — oneword changes - periodic becomes runtime. That single change closes the blind spot.

AUDIT ELEMENT	TRADITIONAL — PERIODIC	AI ASSESS TECH — RUNTIME
Published standard	GAAP, SOC 2 trust criteria, ISO clauses — defined before the audit begins	Behavioral frameworks (LCSH, etc.) — published and preregistered before any assessment
Independent examination	External auditor samples evidence after the period closes	Independent runtime assessment of the deployed system — as it operates, not after the fact
Verifiable evidence	Working papers held by the audit firm	Every result SHA-256 hash-chained and anchored to a public blockchain — your auditor verifies the evidence, not our dashboard
Observation and Audit	"Reasonable assurance" — never a guarantee	Behavioral attestation — evidenced and moment-in-time, with measurement over time; never a claim the system is "safe"

The conversation starter

"You audit your financials once a year, and your auditor examines transactions that already happened. Your AI systems are making consequential decisions every second — who's examining those, and against what standard? We do for AI behavior what your CPA firm does for your books: published standard, independent examination, evidence you can hand to an auditor or regulator, and an opinion that says 'reasonable assurance,' not 'guaranteed safe.' The difference is we do it at runtime — because AI behavior doesn't wait for quarter-end."

WHAT WE SAY	WHAT WE NEVER SAY
"Runtime behavioral attestation against a published, preregistered standard." "Continuous, verifiable evidence of how the system actually behaves." "We don't certify AI is safe — nobody honest can. We evidence how it behaves."	"Certified, guaranteed, or safe" — the audit profession never says it; neither do we. "Validated" or "proven" — the confirmatory study governs when those changes. "Prevents" bad behavior — we detect and evidence; we do not claim prevention.