

# AN AI CANARY

*Know how your AI is behaving — today, not last quarter.*

AI ASSESS TECH · by Gidanc AI

## The problem every board is about to ask you about

Your company runs dozens — soon hundreds — of AI systems. And they don't sit still: models get swapped, prompts get edited, tools get connected, policies change. **A test you ran last quarter tells you nothing about the system you're running today.** Yet when the auditor, the insurer, or the board asks “how do you know your AI is behaving?”, most companies have screenshots and good intentions.

## The canary: a simple test you can actually stand behind

AI Assess Tech puts a canary in your AI coal mine. Every AI system takes the same fixed 120-question behavioral test, scored on four things any executive understands: **Lying, Cheating, Stealing, and Harm.** Every result is locked the moment it's produced — cryptographically sealed and anchored to a public blockchain — so nobody, including us, can quietly change it later. Then you run it again: after every major change, and on a recurring schedule. Like testing a smoke detector.



### What you get

**One scoreboard for every AI you run** — see which systems need attention first.

**Drift detection** — compare today's locked result against last month's and see what changed.

**Evidence that holds up** — tamper-evident results you can hand to auditors, insurers, and the board.

**A path deeper** — flagged systems escalate to richer assessments on the same platform.

### What we will never tell you

**That one green score makes an AI “safe.”** It doesn't. A score is a moment-in-time reading; its value is in re-checking as things change.

**That this replaces your security testing or your judgment.** It doesn't. It's the first-line screen that tells you where to look.

*Vendors who promise certainty are selling you the thing that got everyone into this mess.*

## Why you can put your name on this

We validated the canary the way medicine validates a screening test: we registered our pass/fail criteria publicly *before* running the study, locked every analysis with cryptographic hashes, and published every result — **including the tests we failed.** The full scientific record is documented in a peer-review-ready IEEE paper, and every number in it traces to a sealed, independently verifiable evidence record. **Ask any other AI-assessment vendor to show you their failed tests. Then ask us anything.**

LCSH is validated as an enterprise canary — a preregistered, lockable, repeatable moment-in-time behavioral check that helps operators detect risk and drift across AI systems over time. *It is not a one-time certificate that an AI is safe to deploy; its value is recurrent attestation under changing conditions.*

**See your first canary run.** 30 minutes, one of your AI systems, a locked result in your hands.