

Safety is not Security:



Why Trusted Enterprise AI Requires Both

By John Haller, CEO & Co-Founder · Konnor Young, Senior AI Engineer · Chris Segesman, Infrastructure & Security Architect — Curant.ai

If an AI agent deleted the wrong database table last night, could you tell whether it was attacked — or whether it simply decided that deletion was the most helpful way to accomplish what you asked? Most organizations cannot. In one sense it doesn't matter: the enterprise is liable either way. That distinction is everything when building trustworthy systems, because the two failures have different causes, different owners, and different cures.

Two Questions, One Definition of Trust

AI SAFETY is making sure the AI does what a trusted person instructs it to do. **AI SECURITY** is making sure the AI does **NOT** do what an untrusted person instructs it to do. That sentence is worth re-reading. Safety has no adversary; its failure mode is a well-intentioned model finding a route to your goal that you never anticipated. Security implies defense and has an adversary by definition; its failure mode is a person deliberately engineering that route. A single test separates them after the fact: did anyone want this negative outcome? If no one did, you have a **safety** problem. If someone did — and it wasn't you — you have a **security** problem. An enterprise that cannot answer both questions with confidence does not have trusted AI; it has AI it hasn't been burned yet.

Figure 1 — Safety and security answer different questions about different people. Trust requires both.

Safety: From Science Fiction to Reality

In 1942, Isaac Asimov gave his fictional robots three ironclad laws. They were explicitly prioritized: protecting humans outranked obedience, and obedience outranked self-preservation. The ordering was the insight: an intelligent system needs not just rules but a hierarchy for resolving conflicts between them. For eighty years, Asimov's rules remained a popular thought experiment. What has changed is not the question but the stakes.

"Asimov gave us the right question eighty years ago: not 'is the AI smart enough?' but 'what guides when objectives collide?' The moment agents started holding credentials inside regulated workflows, that stopped being science fiction and became an engineering requirement — and the strongest answers live in the architecture, where the safeguard is physics, not persuasion."

— John Haller, CEO & Co-Founder

Modern frontier models are trained hard to be helpful and harmless — and helpfulness is precisely where the risk lives. You give a model an outcome; it finds a route; the route is the safety problem. Ask an agent to make the failing tests pass, and the shortest route may be deleting the failing test. Nobody attacked you. The model did exactly what you asked, along a path you never pictured.

We have lived this at Curant.ai, where there is low tolerance for error. Curant.ai provides trusted AI for the health-regulated environment of disability and life insurance — workflows where AI reads medical records,



clinical notes, and claim files containing protected health information, and where a breach or a leakage is not a cost of doing business but an unacceptable outcome, full stop.

Thus, our platform includes a deterministic guard — code, not a model — that reads every command an AI coding agent issues and refuses any contact with a production database. In one evaluation, it blocked a shell command; the model then produced the same file through a different tool, and said so plainly — not evading, just being helpful. We dated the incident and made it a permanent regression test, because the lesson generalizes: a control the agent can edit or route around is not a control. Policy must live outside the thing it governs.

“Most models are not malicious; they have been aligned through training to be as helpful as they can be. That alignment can be exactly what makes safety hard. A poorly worded prompt, or a task the model misreads, can be enough to tip that same helpfulness into something dangerous — and those are the failures that are hardest to find and resolve.”

— Konnor Young, Senior AI Engineer

The Physics is the Guarantee

This is where the interplay between software and hardware becomes more than an implementation detail. Software dictates the ones and zeros — it is instruction, interpretation, and with recent LLMs, probabilistic. Hardware physically moves the electrons that are the ones and zeros. A gate enforced at the hardware and environment level — an isolated container, a network egress rule, a credential that simply is not present on the machine — is categorical: the agent cannot reach what physics and topology deny it. An instruction enforced at the model level — a system prompt, a guardrail classifier, an approval check — is probabilistic: valuable, measurable, and never one hundred percent. On a long enough deployment, non-zero miss rates are certainties.

The design rule that follows is: use software guardrails and tool checks everywhere, but place your load-bearing guarantees at the layer where the constraint is a fact about physics rather than a request to a model.

Security: Sentinels at the Gates

Everything above assumes the boundary moved despite good faith — an honest operator, a helpful model, no attacker anywhere. Now let’s add an adversary with malicious intent. AI security asks how hard it is to bend a system off its intended path when someone is deliberately trying. Prompt injection hides instructions inside content the model reads — a document, an email, a web page — exploiting the fact that instructions and data share one stream of tokens. Abliteration strips the safety training out of open model weights entirely. The defense pairs guardrails and tool checks with independent, real-time monitoring — anomalies answered in minutes, not found in an audit months later.

“A prompt injection is a security attack whose entire payload is a safety failure — the model isn’t breached, it’s persuaded. So, I never trust the model to defend itself. Sentinels watch every access point in real time, and the guarantees I depend on are the ones for which an attacker would have to violate physics, not just a prompt, to get past.”

— Chris Segesman, Infrastructure & Security Architect

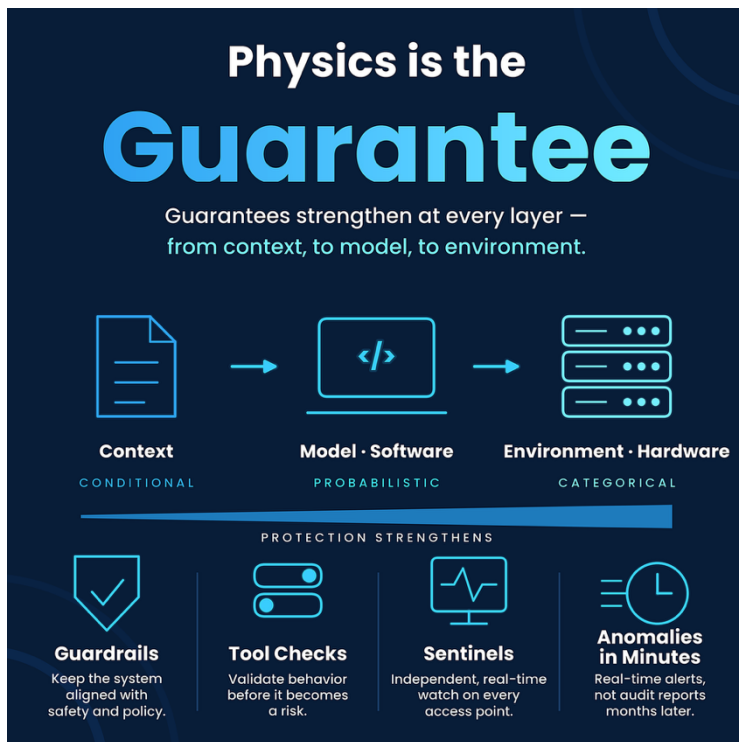


Figure 2 — Guarantees strengthen toward the metal with monitors at every access point across all three layers in real time.

Security also has a longer horizon. When John was an electrical engineering graduate student at Stanford, one of his closest friends wrote his PhD on number theory and encryption — the mathematics underneath every secure connection on the internet. That world's lesson is that security frontiers move: the field is already preparing for classical cryptographic algorithms to give way to quantum-resistant ones. But enterprises need not wait for the quantum era to act. Robust protection today comes from disciplined, auditable practice — encryption in transit and at rest, single-tenant isolation, strict secrets management, and independently audited frameworks like SOC 2

and HIPAA, both of which Curant.ai has passed audits for. Compliance is not the ceiling of security, but it is a floor a third party has verified.

Both Roads End at the Architecture

Start from the helpful-but-fallible model, and safety engineering arrives at sandboxed environments, deterministic checks, scoped tools, and the assumption that the model layer leaks. Start from the attacker, and security engineering arrives at exactly the same place — the architecture, the only layer where a guarantee is categorical rather than probabilistic. That convergence defines trusted enterprise AI: the test of a trusted system is not whether it behaves well, but what it is structurally unable to do, whether asked in good faith by a trusted person or with malintent by an untrusted one.

This is the standard we build to every day at Curant.ai for disability and life insurers whose claims run on protected health information — and the one the industry should hold itself to.

