

# The Live System Model For Agentic Security

## Your board is asking whether your AI security investments are actually working

Your team is making relentless modifications trying to get AI to work consistently across your environment. But it takes longer than expected. They are both looking at the same problem from opposite ends.

Every AI security platform available today runs on telemetry that was never built for AI. Logs, alerts, events from dozens of disconnected tools. That telemetry was designed for someone who already knows the environment: a team member who carries the business context, understands the infrastructure, and fills the gaps. AI gets the events. It gets none of the knowledge.

It is not the AI. The industry has been adding more tools and more AI on top of existing security stacks. But the underlying architecture never changed.

Connecting your data is not the same as understanding your environment. Stream is the only platform that maintains a live model of what your environment actually is, not what your logs recorded.

When AI runs on the wrong foundation, errors do not stay contained. They scale. When the foundation is right, every agent decision compounds the advantage.

"AI triage has really become a new detection layer for us. Investigations that used to take hours now take minutes, and for a small team, that's changed everything about how we operate."



Mike Young  
Director of Cybersecurity, Hunt Energy



## Stream built a different foundation:

CloudTwin is a continuously maintained model of your entire environment: Cloud, on-prem, SaaS, AI Agents in production. It updates the moment anything changes. Every AI agent in your stack runs on that model instead of log patterns.

CloudTwin validates remediation and containment actions before they execute, so AI can act safely without disrupting the business. Not faster guesses. Ground truth.



**AI makes your security faster.  
Stream makes it real.**

Visit to learn more  
[www.stream.security](http://www.stream.security)





# Stream built CloudTwin: a continuously maintained model of your production environment.

Stream ingests the same logs as your SIEM and AI SOC but created a model AI can work with. Instead of enriching alerts, Stream uses that telemetry to continuously update a live graph of your environment: every identity, every permission, every infrastructure

relationship, every attack path, every blast radius. The environment is understood before the investigation begins. Agents verify AI decisions in seconds instead of rebuilding investigations from scratch. That is what creates the trust required for AI to act autonomously.

|               | Legacy context      | AI-native context                 |
|---------------|---------------------|-----------------------------------|
| Form          | Metadata on events  | Live environment model            |
| Updates       | When an event fires | As the environment changes        |
| Built for     | Human reading       | Machine reasoning                 |
| Answers       | What happened       | What the environment is right now |
| Trust ceiling | Human must validate | Agent acts on ground truth        |

### Stream covers your entire attack surface:

- Cloud
- On-prem / Virtualized infrastructure
- SaaS
- AI agents in production

Across every surface, out of the box

- Fuse
- Detect
- Investigate
- Respond
- Scale with AI

## Go further with Stream.Force

Once you see and understand what is happening across your environment, StreamForce lets your team build custom AI security workflows without being an AI expert. Automate what matters most to you, grounded in a live model of your environment, not generic templates.

Alert Triage

Identity Risk Review

Attack Path Analysis

AI Workload Monitoring

Lateral Movement Detection

Compliance Evidence Collection

### Five questions to ask any AI security vendor

1

Does the AI operate on a live model of your environment, or does it reconstruct context from logs every time it acts?

2

When your environment changes, does the AI know immediately, or does it find out when the next alert fires?

3

Can the platform detect threats in AI agents and workloads already running in your production environment?

4

Does a single model cover your entire environment: cloud, identity, on-prem, SaaS, and AI systems, without gaps?

5

Can the AI assess blast radius and consequences before taking an autonomous action?

