

Meet Vector

Autonomous red team, **validated against the controls you already own.**

CTEM is incomplete without AEV. AEV is incomplete without closure. Most exposure validation proves what is exploitable and stops there. Tuskira proves it against the controls you have actually deployed, then closes the path through the tools you already own and re-tests to confirm it stayed shut. **Vector is the outside-in entry point to that loop, and it is available now.**

WHAT VECTOR DOES

Vector probes your approved external scope the way an attacker would, using current TTPs, newly disclosed vulnerabilities and AI-driven techniques. Then it answers the five questions a scanner cannot.

- **Is it exposed?**
- **Can an attacker use it?**
- **Do existing controls already stop it?**
- **If not, where does it lead?**
- **What is the fastest way to close the path?**

Tested, not guessed.

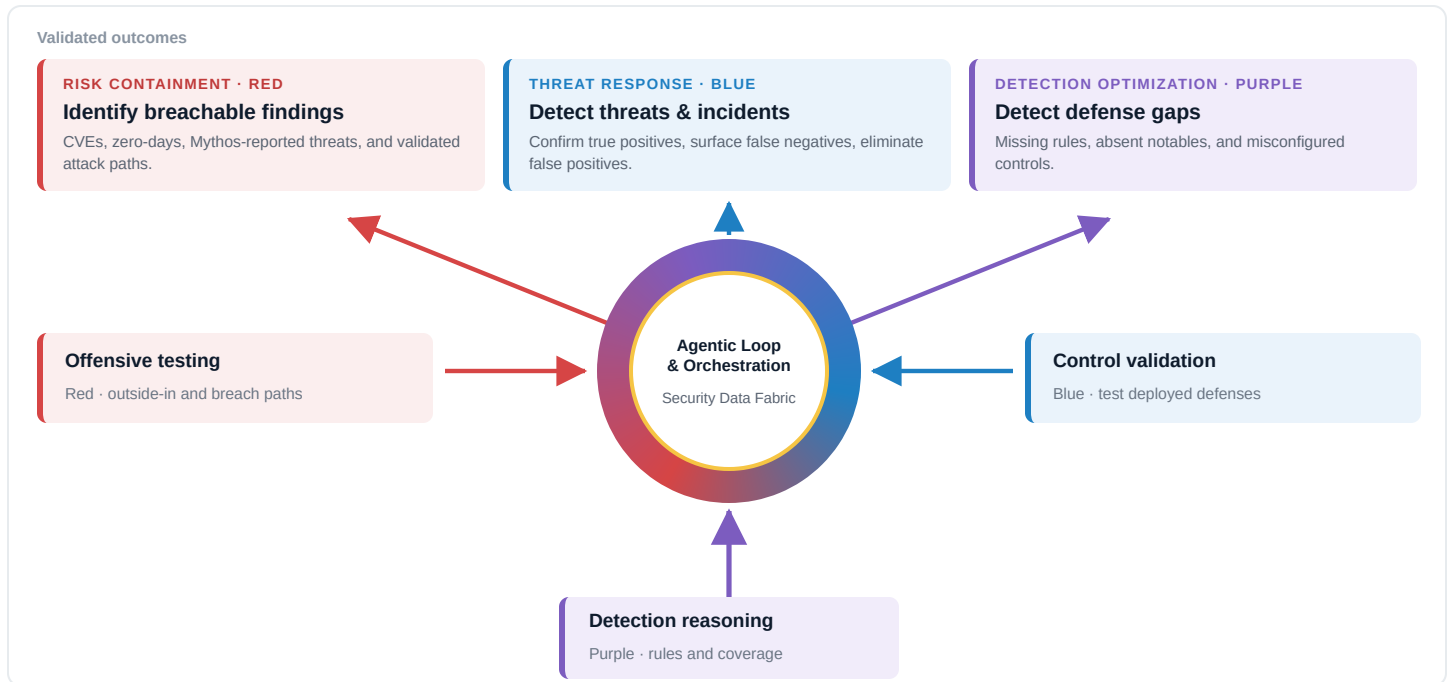
A finding your WAF already blocks never becomes work. A finding nothing stops arrives with the attack path, the blast radius and the control change that closes it.

Runs continuously inside customer-approved scope, grounded in your own architecture rather than blind exploitation against production.

THE ARCHITECTURE

One agentic core, three validated outcomes

Offensive testing, control validation and detection reasoning run over one Security Data Fabric, continuously producing the findings, threat calls and defense-gap fixes each team needs. Vector, Kairo, Lattice, Quell and Iris all work on that same model.



HOW THE LOOP RUNS

01 SENSE

Attack it

Vector probes your approved external scope the way an attacker would.

02 NORMALIZE

Model it

Security Data Fabric, one neutral layer for every signal.

03 VALIDATE

Prove it

Validate whether your deployed controls block the path.

04 CLOSE

Shut it

Close the path through the controls you own, tune detection.

05 RE-TEST

Confirm it

Re-run the path to prove it's actually closed.

VECTOR AND THE REST OF THE LOOP five agents, one model of your environment

Vector NEW	Can they get in?	Outside-in red team agent. Probes your approved external scope with current TTPs, then validates each finding against your deployed controls.
Kairo	Where can they go?	Maps and validates the cross-domain breach path to the crown jewel, including lateral movement across cloud, identity and on-prem.
Lattice	What matters most?	Correlates internal risk and decides what is truly exploitable and worth fixing first.
Quell	Does this CVE open a path?	Answers reachability for a newly disclosed vulnerability, here, now.
Iris	What happened, what next?	Investigates alerts with asset, identity and blast-radius context, then stages response under policy.

Vector answers the first question. Everything after it runs on the same model, so a Vector finding arrives at the next agent with its asset, identity and control context already attached.

WHY VECTOR IS DIFFERENT

It checks your controls

Most external testing stops at reachable. Vector asks whether your **deployed EDR, WAF, NGFW and IAM** already break the path.

It knows your architecture

Findings land in a model of your cloud, identity, endpoint, network and SIEM, so Vector validates **without blind exploitation against production.**

It never stops running

Not an engagement with an end date. When the environment changes, **Vector re-tests and the answer stays current.**

Governed by design: testing and analysis run autonomously. Containment requires approval. Policy decides whether a response is recommended, staged or executed, and execution runs through the controls you already own.

Your data stays yours: logs stay where they live and are reached through federated search. Every hypothesis, tool call and verdict is logged per tenant with the full evidence chain. Nothing trains across customers.

WHAT VECTOR LOOKS FOR

EXTERNAL SURFACE

What attackers see

Internet-facing assets, services, identities and misconfigurations, in approved scope.

EXPOSURE

What is reachable

Call-graph reachability and breach paths across cloud, identity, endpoint and on-prem.

CONTROLS

What already stops it

Whether your deployed WAF, EDR, NGFW and IAM break the chain today.



THE OUTPUT

A closed path

Proven shut, with evidence

THE PROOF, IN NUMBERS

The whole point of the loop is to spend the team's time only on what is real. In Tuskira deployments, Kairo has deprioritized **up to 99% of scanner findings as unreachable**, whichever tool surfaced them, then proves the survivors against your live controls and closes them. **Less noise, less patching, fewer tools to run the same job.**

99%

of scanner findings deprioritized as unreachable, so teams work only what is exploitable. · Source: Tuskira, Kairo launch.

QUESTIONS WE GET

"WE ALREADY HAVE A CNAPP."

Good. It finds cloud and code exposure, and Tuskira reads it as one more signal. **Neither a CNAPP nor a scanner tests whether your deployed controls stop the path** it just flagged.

"WE ALREADY RUN PENTESTS AND BAS."

Those prove a path was exploitable on the day they ran. **Tuskira runs the same validation continuously against current state**, so the answer does not expire with the report.

"OUR PLATFORM VENDOR IS BUILDING AGENTS."

Most of those loops reason over that vendor's own telemetry. **Real environments are multi-vendor, so the model of that environment has to be multi-vendor too.**

SEE IT AGAINST YOUR OWN ENVIRONMENT

Most exposure validation proves a path is reachable.

Tuskira runs the full loop. Prove the path, validate the controls, close it through the tools you already own, and re-test to prove it stayed shut. One neutral model across every tool you already run.

REQUEST A DEMO

tuskira.ai/request-a-demo

We will run Vector against your approved external scope and show you what your controls already stop.