

REALM + VEGA

Agentic defense is only as strong as the data it reads

Realm is the SOC-aware security data pipeline: full control of your telemetry between source and destination. Vega's Agentic Cyber Defense Platform runs detection, triage, and analytics on your data wherever it sits, with no migration, ingestion, or egress. Realm sits upstream: it reduces, enriches, redacts, and normalizes your telemetry on its way to the systems Vega reads. The result is agentic defense that reaches verdicts faster, hands back fewer inconclusive investigations, and works at its actual ceiling.

The center of gravity in the SOC has moved

Detection, triage, and analytics are moving to where the data sits. That move does not settle the data question: which sources were collected, what condition they landed in, and what history stays reachable.

Most pipelines cannot answer it. They are transport layers: rules written once, a little filtering, hand-tuning when something breaks. That held up when there was one destination and requirements changed once a year.

It does not hold up now. Workflows draw on the SIEM, cloud platforms, and storage at once, and what each one needs from the data changes faster than anyone can hand-tune.

Vega can normalize what it finds. What it cannot do is recover a field that was never collected, or reach a source that never landed. That is decided upstream: which sources were collected, and what they contained when they landed.

So the team has to run a pipeline that understands the data and makes decisions about it, and know the systems Vega reads are complete before an investigation finds the gap.

What changes with Realm upstream

THE STATUS QUO

Static filters written once decide what every destination sees. Volume drops, and nobody can say what was cut until an investigation goes looking. Every new destination means re-plumbing collection, source by source.

WITH REALM

Realm reads what each source contains, field by field, and what each destination needs, then recommends what to keep, cut, normalize, enrich, and route, with the reasoning shown. Your team approves, and everything not sent downstream is retained in Data Haven, searchable and ready to be resupplied. **Collect once, cut with proof, keep everything.**

AT A GLANCE

Independent collection

Add or change destinations without re-plumbing a single source.

In-stream enrichment

IP intelligence and threat context attached before data lands anywhere.

Detection Integrity

Every reduction checked against the detections you run, mapped to MITRE ATT&CK.

Complete retention

Everything not sent downstream is retained in Data Haven, searchable directly.

Data composition

When volume moves, see the exact field and value driving the change.

0

Detections lost at Vensure while SIEM volume fell 83%.

What Realm does for Vega

The full picture, protected

- **Detection Integrity:** Before any reduction ships, Realm checks it against the detections you run and protects the fields they depend on, mapped to MITRE ATT&CK. Nothing Vega works from goes silently missing.

Context for faster verdicts

- **Enrichment:** Realm attaches IP intelligence and commercial threat intelligence while data is still in stream. Events arrive carrying the context an agent would otherwise assemble mid-investigation, so verdicts come faster and fewer end inconclusive.

AI-ready data

- **Realm Lake format:** Every event lands in your lake as one consistent record that holds the raw log, the parsed fields, the observables, and the enrichments. There is no normalization project first. Your team can put an agent on it the week it lands, not six months later. The format is open JSON and destination-agnostic: proven in production today, and the same record works wherever your AI reads from next.

Key outcomes

- **Lower cost, more context:** Cut SIEM ingestion while Vega works from enriched, complete, consistent sources.
- **Detection coverage held:** Reductions are validated against the detections you run before they ship.
- **Better verdicts, fewer dead ends:** with more context in every event and full history in reach, more investigations resolve instead of escalating back to your team.
- **Vendor independence:** Change SIEMs, storage, or analytics tools, even the SIEM itself, without re-plumbing collection.

If AI is going to run the SOC, AI should be preparing the data it runs on

Vega runs on top of the systems you already have. It does not sit at the source: it does not decide which sources are collected, or what they contain when they land. It reasons over what exists, and that is the right design.

Realm is where your team governs what exists: what reaches the systems Vega works from, what context it carries, which detections stay protected, what stays reachable. Only Realm checks every reduction against the detections you actually run before anything ships. Realm prepares the context. Vega makes the call, with a higher ceiling on what it can find. That is the whole idea behind Realm.

GLOBAL MANUFACTURER

One of the largest Microsoft Sentinel deployments in the world.

3×

return per \$1 spent

1,000+

KQL rules replaced

72 hrs

to replace them

Prove it on your data.

Pick one source and start a 7-Day Data Assessment. We show the reduction and the detection integrity, on your data.

[Start an assessment → realm.security](https://realm.security)

ABOUT REALM SECURITY

Realm Security gives security teams independent control over their telemetry before it reaches the SIEM, data lake, or AI agent. The platform understands what each source produces and what each downstream system needs, then guides teams on what to retain, cut, route. Headquartered in Boston, Realm deploys in about a week and typically cuts SIEM ingestion costs by 50% or more.



www.realm.security

LinkedIn: /company/realm-security

X: @Realm_Security