

Cisco Secure AI Factory Case Study

Reference Security Design for Nvidia-based
Retrieval Augmented Generation

July 2026

Introduction

This Cisco Secure AI Factory case study is for network architects, security engineers, platform and SRE teams, and AI-platform owners evaluating secure on-premises AI. It describes the lab-validated Atlanta build. It also shows how each layer is configured, what was validated, and what remains on the roadmap.

Each layer section follows the same structure, and each subsection answers one question about the layer:

- **Role** – why the layer exists and what it is accountable for. It names the layer's single responsibility, where it sits in the seven-layer and trust-boundary models, the boundary of its concern (what it deliberately does not do), and what must exist below it. Read the Role subsections top to bottom and they form a standalone narrative of the whole design, with no product names or configuration.
- **Components** – what the layer is made of. The products, versions, and workloads that implement the role.
- **Design and implementation** – how it is built. Addressing, topology, policy, configuration, and – where useful – a numbered implementation overview.
- **Operational and validation state** – what was deployed and proven. The evidence, with scope notes where coverage is intentionally bounded.
- **Design considerations** – what to watch when building it. Grounded constraints and deployment lessons, written as actionable takeaways.
- **Dependency handoff** – what the layer hands to the layers above it, so each section can be read standalone while its place in the stack stays clear.

Note: Validated capabilities and roadmap items are stated plainly and separately. Where the lab surfaced a constraint, a gap, or an unresolved defect, this document records it rather than omitting it – a validated design is only useful if its boundaries are legible.

Executive summary

The Atlanta Secure AI Factory (SAIF) is a validated, on-premises AI platform for enterprise AI workloads, including retrieval-augmented generation (RAG). Security and observability are built into each layer instead of added only at the perimeter. The build uses two Red Hat OpenShift clusters on Cisco UCS compute and a Cisco Nexus 9000 VXLAN-EVPN fabric. It runs the NVIDIA AI Enterprise runtime (NIM, NeMo, Milvus), uses Isovalent Networking for Kubernetes (Cilium) and Isovalent Runtime Security (Tetragon) for workload-layer control, uses Cisco AI Defense at the model layer, and routes signals through a Splunk OpenTelemetry Collector into a two-plane Splunk analytics tier.

The two clusters have distinct roles:

- The **AI Defense cluster** runs the Cisco AI Defense on-premises data plane and its supporting inference models.
- The **AI RAG Application cluster** runs the NVIDIA inference, retrieval, and vector-store workloads that make up the RAG pipeline.

Separating the clusters keeps AI Defense classifier models off the GPUs that serve user inference. It also turns cross-cluster traffic into an inspection point rather than an internal detail.

Results at a glance. The highlights below are drawn from the validated build; each is detailed, with its source, in the sections cited.

- **Layered controls are deployed.** The stack maps 5 nested trust boundaries across 7 enforce-and-observe layers. In this build, B3-B5 are deployed; B1-B2 remain roadmap. Identity-based policy and kernel-level runtime observation reduce blast radius inside the deployed boundaries.
- **Model protection runs on-premises.** Cisco AI Defense guardrails for prompt injection and malicious responses were validated on the on-premises data plane through the Inspection API and Gateway (TC-03, TC-05, TC-06). This validates the runtime inspection path; independent egress-capture proof is not part of this phase.
- **Runtime security observed:** a sustained ~130,000–160,000 Tetragon/Hubble runtime events per day streamed into Splunk from the workload cluster.
- Security analytics, firing end-to-end: Risk events feed Risk-Based Alerting, and Enterprise Security Findings are generated from live SAIF signals once RBA thresholds are met.
- **Platform footprint: the build runs on 2 Red Hat OpenShift clusters across 10 Cisco UCS nodes.** The node set is 6 C225 M8S control-plane nodes and 4 C845A M8 GPU workers. Those workers host 16 NVIDIA RTX Pro 6000 GPUs serving five models. The fabric is Cisco Nexus 9000 VXLAN-EVPN with a 200 GbE data plane; see Table 4 for the full BOM.

Business and customer outcomes

Enterprises adopt on-premises AI when they need to use proprietary data, keep inspection points auditable, and give security teams visibility into AI-specific risks. This design addresses those needs by placing AI runtime, model protection, workload security, telemetry, and analytics inside a layered enterprise control model. The table below states each customer outcome with the mechanism used in the lab and the current validation state.

Table 1. Customer outcomes, proof mechanism, and validation state

Customer outcome	Mechanism and lab evidence	Validation state / gap
Keep sensitive AI workloads on enterprise-controlled infrastructure	AI Defense Hybrid Connector; on-premises AI Defense data plane; on-premises NVIDIA inference/RAG workloads. TC-03, TC-05, and TC-06 passed against the on-premises Inspection of API and Gateway paths.	Validated for the runtime inspection path. Independent egress capture proving no prompt/response egress is not included in this phase.
Reduce blast radius with layered controls	B1-B5 trust-boundary model; B3-B5 deployed with OpenShift, Cilium, Tetragon, and AI Defense model-interaction control. Runtime and flow telemetry are observed in Splunk.	Partial. B1-B2 are roadmap; default-deny negative tests and failure-injection tests remain pending.
Give SRE and SecOps shared operational/security visibility	Splunk OTel Collector, Observability Cloud, Splunk Enterprise + ES, and shared Kubernetes identity through k8sattributes. Two-plane telemetry is deployed; AI Defense events reach ES; Findings fire from live SAIF signals.	Validated for deployed paths. Phrase as shared identity across purpose-built telemetry paths, not one telemetry pipe.

Validated scope and roadmap

This is a lab-validated design, not a fully hardened production deployment. The distinction is stated explicitly, so the reader can rely on the validated path and plan for the rest. The customer outcomes in Business and customer outcomes are stated with their validation state; this section defines the proof boundary so validated, partial, and roadmap items are not conflated.

Validated in this build:

- VXLAN-EVPN fabric segmenting the two clusters into separate overlay networks; Layer 3 handoff to the edge.
- OpenShift clusters on UCS with NVIDIA GPUs; the NVIDIA AI runtime (NIM/NeMo/Milvus) deployed.
- Cisco AI Defense runtime inspection on both consumption surfaces: Single-turn red-team validation (TC-01), prompt-injection detection via the Inspection API (TC-03) and the Gateway (TC-05), and malicious-response detection via the Gateway (TC-06).
- The Splunk two-plane telemetry pipeline; AI Defense events integrated into Splunk Enterprise Security through the Cisco Security Cloud App.

On the roadmap:

These are recorded in the relevant sections and are not part of the validated path.

- The **north-south firewall** (Cisco Secure Firewall 4245) is racked but not configured; trust boundaries B1 and B2 are not yet enforced. See *Network and fabric design*- [Design considerations](#).
- **Multi-turn adaptive red-teaming** (TC-02) is not available on the Hybrid Connector deployment shape. See *LLM and model security*- [Design considerations](#).

Solution architecture overview

The architecture rests on five models defined here – the seven-layer stack, the five trust boundaries, the two-cluster topology, the bill of materials, and the load-bearing design decisions – that the rest of the document builds on. Terminology introduced here (layers, boundaries, cluster names, decision IDs) is used without redefinition in later sections.

Seven-layer model

The stack is organized into seven layers. Each layer runs specific products, enforces one class of control, and observes one class of signal. The L1-L7 labels are document-local SAIF layer identifiers, not OSI layer numbers. The principle is consistent: **every layer both enforces and observes**. Enforcement without observation cannot be audited; observation without enforcement is not a control.

Table 2. The seven-layer enforce-and-observe model

L	Layer	Runs	Enforces	Observes
L7	AI Application	Customer code (chatbot / RAG client / agent) in its own namespace	App logic, user authentication, session and tenant scoping	Application logs and (roadmap) OTEL traces
L6	AI Runtime	NVIDIA NIM, NeMo Retriever, Milvus, GPU Operator	GPU scheduling, model lifecycle (NIMService/NIMCache), OpenAI-compatible inference contract	NIM /metrics, DCGM per-GPU telemetry
L5	AI Protection	Cisco AI Defense (on-premises data plane + SCC control plane)	Prompt-injection, sensitive-info-disclosure, jailbreak, and DLP guardrails on prompts and responses	Per-prompt/response verdict events
L4	Runtime Security	Isovalent Networking for Kubernetes (Cilium), Hubble	Identity-based network policy; kernel-level allow/deny on syscalls and file/network operations	Every flow (L3-L7); process exec and matched syscalls per node

L	Layer	Runs	Enforces	Observes
		Timescape, Isovalent Runtime Security (Tetragon)		
L3	Kubernetes	Red Hat OpenShift (RHCOS, SELinux, etc., API server, admission)	RBAC, Security Context Constraints, namespace boundary, admission policy	API-server audit log, kubelet logs
L2	Compute	Cisco UCS C225 M8S (control plane), UCS C845A M8 (GPU workers), NVIDIA GPUs, BlueField-3	Firmware/BIOS policy, server profiles, hardware attestation	Per-node CPU/memory/power/thermal
L1	Network	Cisco Nexus 9000 leaf-spine (VXLAN-EVPN), Cisco Secure Firewall (roadmap)	North-south access; VRF/VLAN traffic-plane separation	Nexus streaming telemetry

Note: In this build, dependencies run in order: L4 provides runtime ground truth, then L6, then L5. The document is organized top-down by layer so readers can use it as a reference.

Five trust boundaries

Layers describe where components run. Trust boundaries describe where the design changes its assumption about who is trustworthy and enforces a different control. The design names five nested boundaries. A compromised pod must cross all five to exfiltrate data to the internet.

Cisco AI Defense is treated as a model-interaction control inside B5, not as a new numbered boundary; B1-B5 remain the inherited infrastructure boundary set.

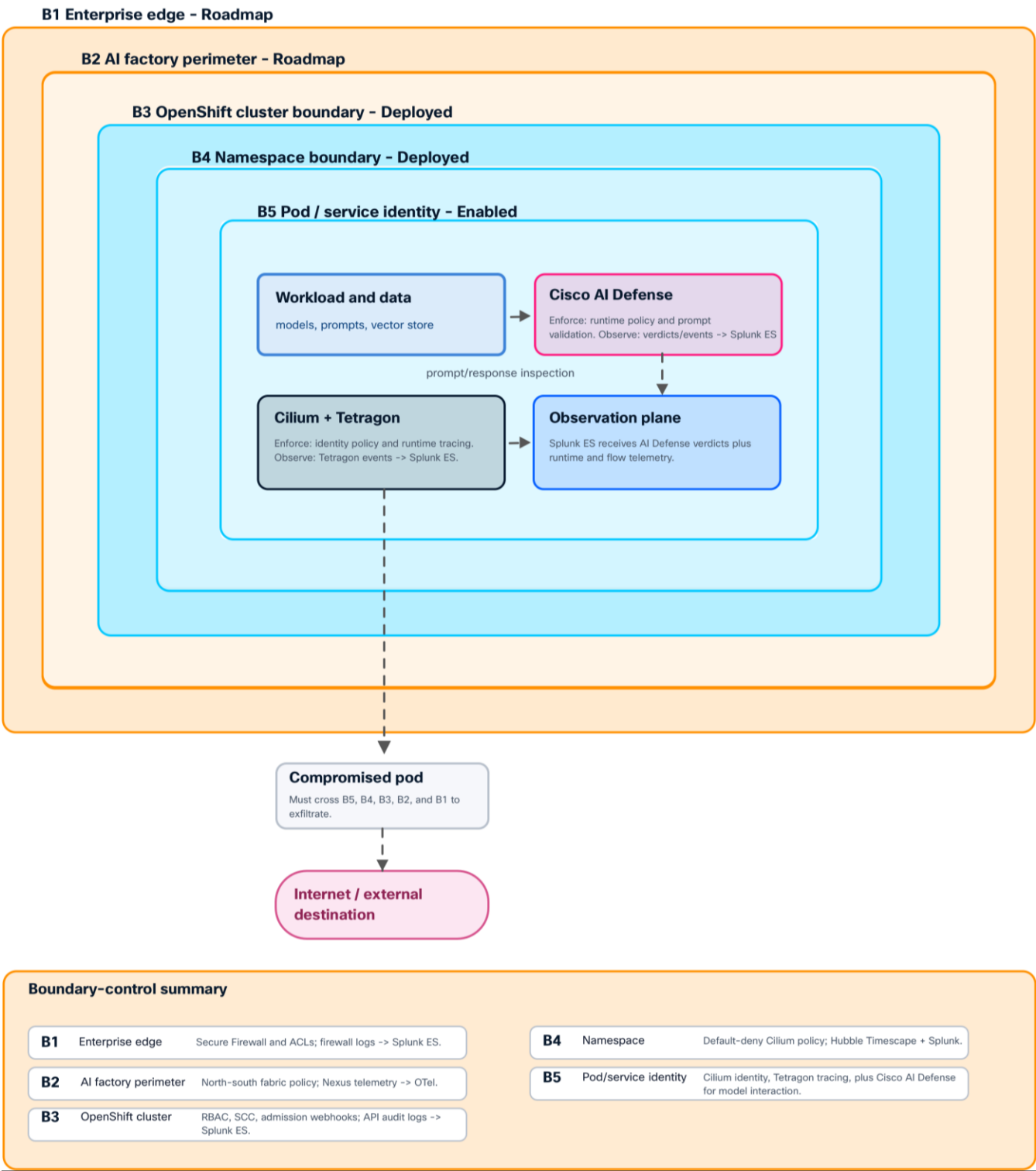
Table 3. Five nested trust boundaries

B#	Boundary	Enforcement control	Observation control	State in this build
B1	Enterprise edge	Cisco Secure Firewall (NGFW), ACLs	Firewall logs > Splunk ES	Roadmap – firewall racked, not configured.
B2	AI factory perimeter	North-south fabric policy at the border-leaf / edge	Nexus streaming telemetry > OTel	Standalone
B3	OpenShift cluster boundary	OpenShift RBAC, Security Context Constraints, admission webhooks	API-server audit log > Splunk ES	Deployed
B4	Namespace boundary	CiliumNetworkPolicy / CiliumClusterwideNetworkPolicy, default-deny	Hubble flow logs > Timescape + Splunk	Deployed; default-deny negative/lateral-movement tests not performed.
B5	Pod / service identity	Cilium identity (label-derived); Tetragon TracingPolicy in the kernel; Cisco AI Defense model-interaction guardrails on prompt and response paths.	Tetragon events > Splunk ES; Cisco AI Defense verdict events > Splunk ES.	Enabled on both clusters. AI Defense runtime inspection validated through the on-premises Inspection API and Gateway (TC-03, TC-05, TC-06); verdict events reach Splunk ES. Runtime events reach Splunk from atl-ocp2; atl-ocp1 runtime-log export remains scoped for Splunk capacity; full standing policy set being confirmed on-cluster.

The five boundaries are nested rather than parallel. Each one sits inside the last, so the data at the center, including models, prompts, and the vector store, is protected by every outer boundary. The diagram below shows this structure, the enforcement and observation controls at each layer, and which boundaries are deployed today versus planned on the roadmap.

Figure 1. Nested trust boundaries B1-B5

SAIF view with Cisco AI Defense represented as the model-interaction control.



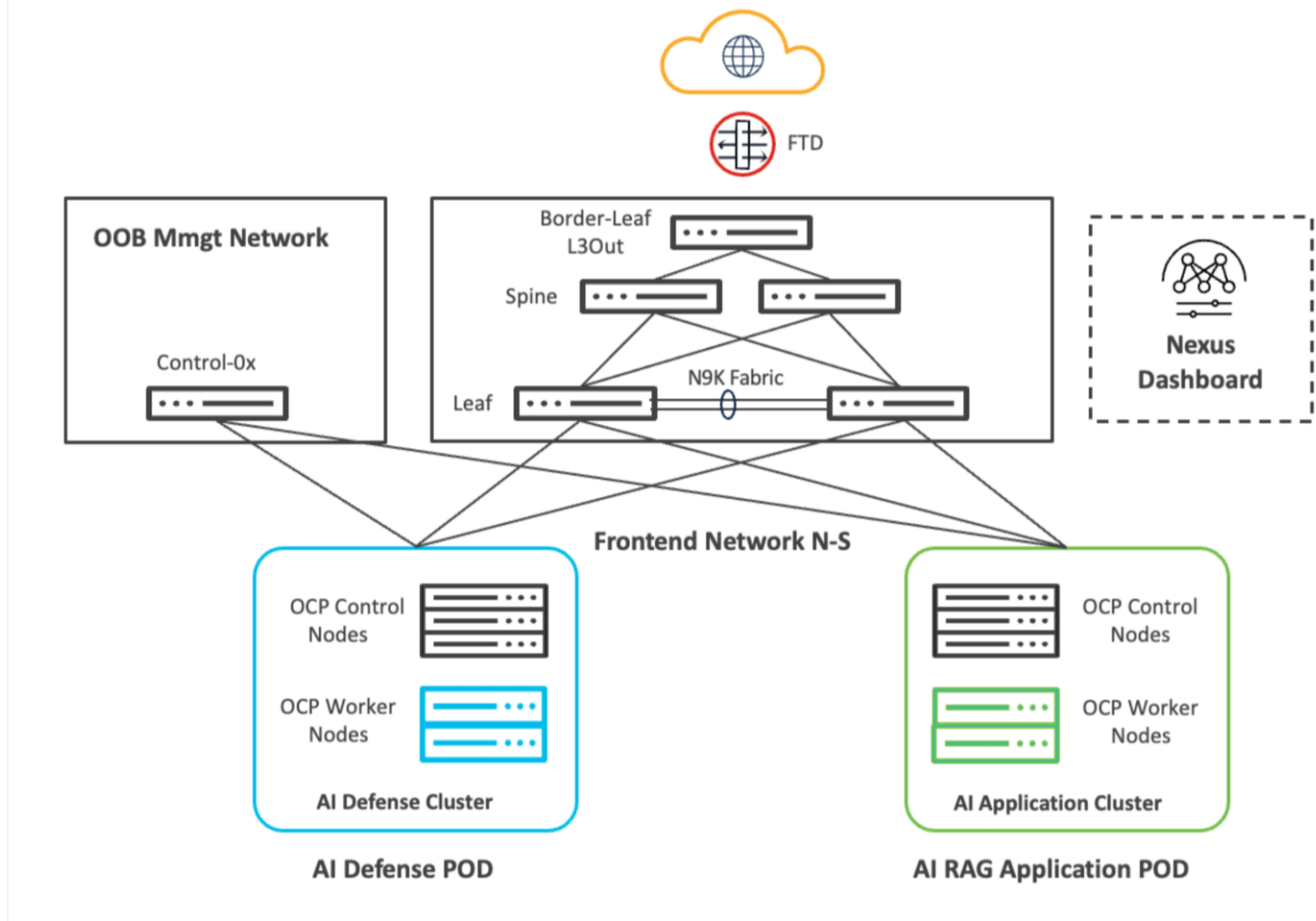
Each boundary changes the design’s trust assumptions and adds its own enforcement and observation controls. B3-B5, covering the OpenShift cluster, namespace, and pod/service identity, are deployed. B1 -

B2, covering the enterprise edge and AI factory perimeter, are on the roadmap. Because the boundaries are nested, a compromised pod would need to cross all five to exfiltrate data externally.

Physical and logical topology

The Atlanta build seeds new hardware for the AI clusters and reuses existing Cisco Validated infrastructure for the fabric and edge.

Figure 2. SAIF Atlanta physical topology

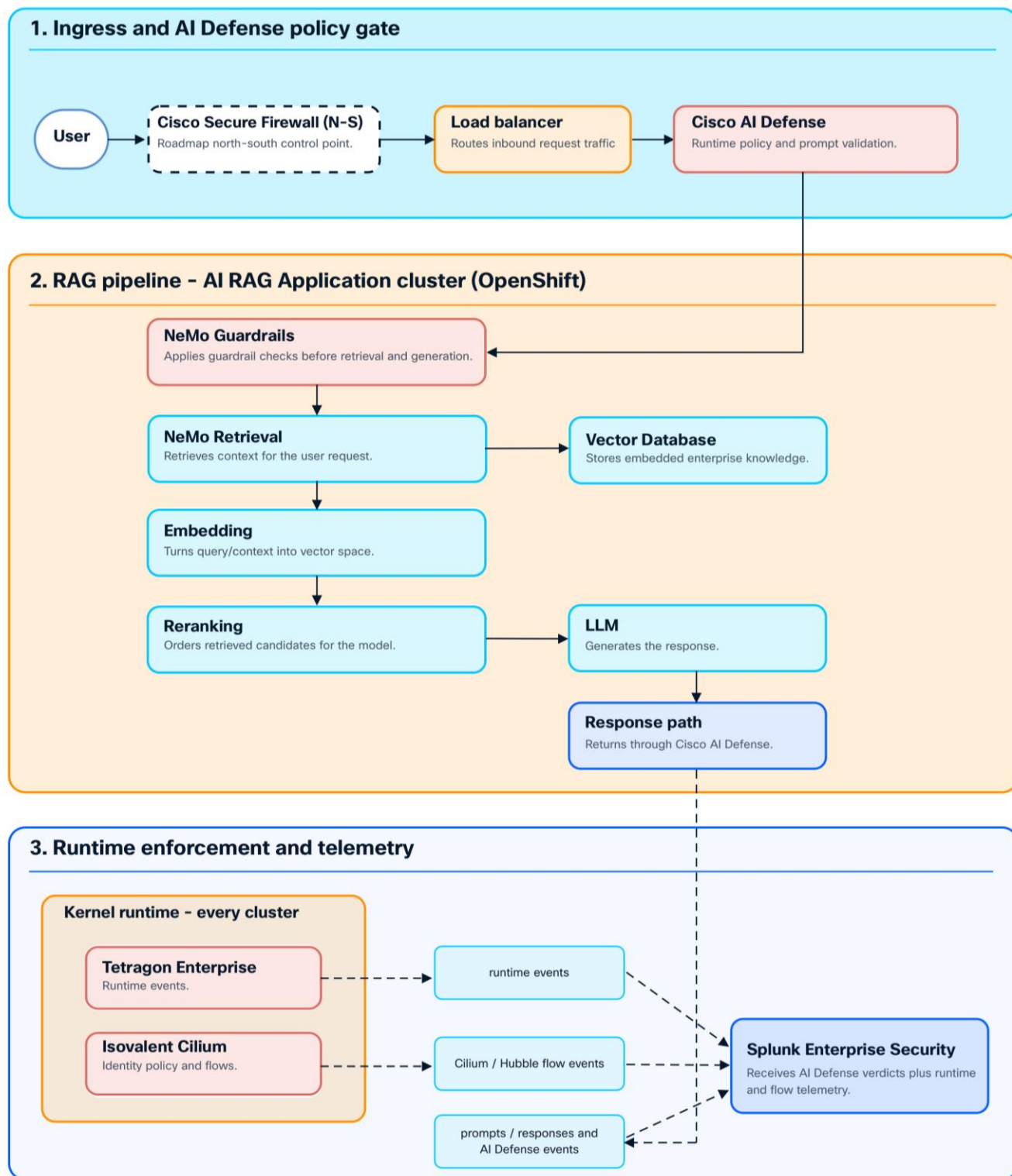


This diagram depicts the two OpenShift clusters (AI Defense and AI RAG Application), each with three UCS C225 M8S control-plane nodes and two UCS C845A M8 GPU worker nodes, dual-homed through NVIDIA BlueField-3 interfaces to the Nexus 9000 leaf pair; out-of-band management on 192.168.0.0/24.

The logical view shows how a RAG request traverses the security and runtime layers.

Figure 3. Logical RAG flow with Cisco security controls

Portrait view of inference, policy inspection, runtime enforcement, and telemetry.



A user request passes the (roadmap) north-south firewall and load balancer, is inspected by Cisco AI Defense, then flows through NeMo guardrails and retrieval (Milvus vector database), embedding, reranking, and the LLM. The response returns through AI Defense. In every cluster, Tetragon observes the kernel and

Cilium enforces identity policy. Telemetry, made up of prompts, responses, AI Defense verdicts, Cilium/Hubble flows, and Tetragon events flows to Splunk. “Guardrails (Future)” markers denote planned, not-yet-deployed points.

Bill of materials

Table 4. Bill of materials (validated deployment)

Category	Item	Quantity	Role
Compute – control	Cisco UCS C225 M8S	6	OpenShift control-plane nodes (3 per cluster)
Compute – worker	Cisco UCS C845A M8 with NVIDIA RTX Pro 6000 GPUs	4	OpenShift GPU worker nodes (2 per cluster)
Network – Ext Fabric	Cisco Nexus N9K-C9332D-GX2B	1	External Handoff/ToR
Network – EVPN VXLAN fabric	Cisco Nexus N9K-C9332D-GX2B	5	Spine / leaf / border-leaf
Network – management	4-node Cisco Nexus Dashboard cluster	1	Fabric management (NDFC); 3+1 node active in this build
Security – perimeter	Cisco Secure Firewall 4245 (“SAIF-FW”)	1	North-south firewall (staged; configuration outstanding)
Host networking	NVIDIA BlueField-3 B3220 DPU (2×200 G)	per server	North-south 200 GbE data plane
Host networking	NVIDIA BlueField-3 B3140H SuperNIC (1×400 G)	4 per worker	East-west fabric (present, not validated)
Host networking	NVIDIA BlueField-3 B3220L SuperNIC (2×200 G)	per control node	Control-node data plane

Software stack

Product	On-prem/SaaS	Version
Cisco Security Cloud Control (SCC)	SaaS	Cisco-managed service
Cisco Secure Firewall (FTD)	On-prem	Staged not configured in Phase 1
Cisco AI Defense	Hybrid	Hybrid Connector
Isovalent Networking for Kubernetes (Cilium) + Hubble Timescape + Isovalent Runtime Security (Tetragon)	On-prem	Operator-managed
Splunk Enterprise + Splunk Enterprise Security (ES)	On-prem	Splunk Enterprise 10.2.1 Enterprise Security 8.4.0

Splunk Observability Cloud	SaaS	SaaS realm us1
NVIDIA AI Enterprise (NIM/NeMo)	On-prem	NIM image example 1.5.0
Red Hat OpenShift Container Platform	On-prem	4.19 (Kubernetes 1.32.x)

Design decisions

The load-bearing decisions for this build are each represented in the section that implements it.

Table 5. Key design decisions

ID	Decision	Rationale	Alternatives/Constraints
D-1	Cilium is the CNI for both clusters, replacing OpenShift's default OVN-Kubernetes	Identity-based policy and L7/FQDN enforcement are prerequisites for the trust-boundary model (B4/B5)	OVN-Kubernetes offers only L3/L4 IP-based NetworkPolicy. See Application and workload security .
D-2	AI Defense runs as Hybrid Connector: on-premises data plane, SaaS (SCC) control plane	Data residency – prompts and responses stay on-premises; only policy/metadata leave	SaaS-managed runtime sends prompts to Cisco cloud; constraint: no multi-turn validation. See <i>LLM and model security</i> - Design considerations .
D-3	The compute data plane is 200 GbE Ethernet via BlueField-3; RoCEv2 is not used	No Backend network i.e. GPU to GPU connectivity desired.	The 4×400 G east-west SuperNIC fabric is present but not used. See <i>Compute, GPU, Storage</i> - East-west SuperNIC fabric .
D-4	The OTel Collector egress allowlist is a closed set of regional SaaS endpoints	Cost control and egress governance	New vendor endpoints require an allowlist update; treat as a versioned artifact. See <i>Telemetry architecture</i> - Dependency handoff .
D-5	Default-deny CiliumNetworkPolicy is the baseline in every AI namespace	Zero-trust posture at B4; required paths added explicitly	Coverage is asserted; negative/lateral-movement tests were not performed. See <i>Application and workload security</i> - Design considerations .
D-6	AI Defense integrates at the shared LLM gateway (integration Pattern 3)	One enforcement decision for all applications, zero application code change	Patterns 1 and 2 remain available for apps that cannot use the shared gateway. See <i>LLM and model security</i> - Design and implantation .
D-7	Service LoadBalancer VIPs are advertised into the fabric by Cilium BGP (eBGP-multihop to the leaves), not by a static external load balancer	Service exposure follows workload identity and state – a VIP is on the wire only while its Service runs and Cilium advertises it; no static fabric ACL to maintain	Requires per-cluster eBGP peering to the leaves (ASN 65500/65501 <> 65525); the pool prefix is learned dynamically, not statically originated. See LoadBalancer IPAM and BGP advertisement .

Dependency handoff

The architecture model, boundary model, and decision register defined here are referenced by every subsequent section. The build order follows the layer stack from the bottom up: the [Network layer](#) must be

operational before the [Compute and OpenShift platform](#), which must be in place before the [security](#), [runtime](#), [telemetry](#), and [analytics](#) layers.

Network and fabric design

Role

The Network layer (L1) provides the switched fabric that carries the data plane for the two OpenShift clusters and connects them to the north-south edge. It establishes the first segmentation boundary: the two clusters do not share a broadcast domain, subnet, or VLAN. Each cluster's data plane is a separate overlay network on a VXLAN-EVPN fabric, so the identity-based controls in later sections operate on top of Layer 2 and Layer 3 separation rather than a shared flat network.

Components

- **Cisco Nexus Dashboard controller** – A 3+1 node Cluster version 4.1(1g), manages both the fabrics.
- **Cisco Nexus N9K-C9332D-GX2B switches** (32 × 400G), NX-OS 10.5(1) – two spines, two leaves, one border-leaf, and one external ToR.
- **NVIDIA BlueField-3 B3220 DPU** per host, dual-homed to the leaf pair; 200 GbE Ethernet north-south data plane (D-3; no RoCEv2).
- **Cisco Secure Firewall 4245** (“SAIF-FW”) – racked, not configured.

Design and implementation

The controller manages two fabrics: a production VXLAN-EVPN fabric and a single-switch external fabric for the Layer 3 handoff.

Table 6. VXLAN-EVPN fabric parameters

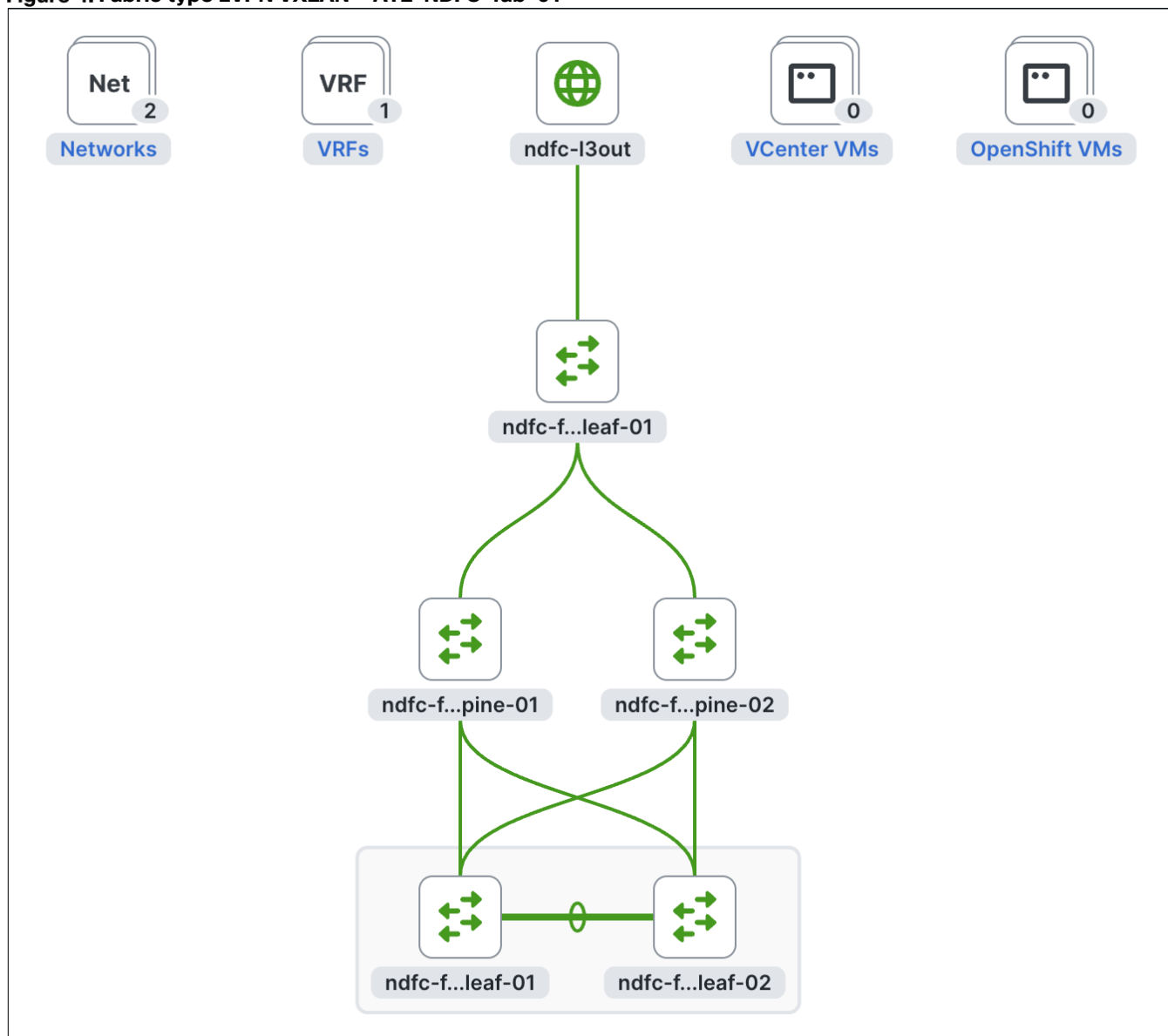
Fabric	Type	ASN	Switches	BUM replication	MTU
ATL-NDFC-fab-01	Easy_Fabric (VXLAN-EVPN)	65525	2 spine (anycast RP), 2 leaf (vPC pair), 1 border-leaf	Multicast, group 239.1.1.0/25	9216
ndfc-l3out	External_Fabric	65522	1 external ToR	n/a	9216

Table 7. Nexus 9000 switch roles

Role	Model	NX-OS	Notes
Spine × 2	N9K-C9332D-GX2B	10.5(1)	Anycast rendezvous points
Leaf × 2	N9K-C9332D-GX2B	10.5(1)	vPC pair; host-facing ports E1/11-E1/20
Border-leaf	N9K-C9332D-GX2B	10.5(1)	eBGP DCI handoff on Eth1/32
External ToR	N9K-C9332D-GX2B	10.5(1)	ASN 65522 termination

Host attachment is split by cluster: the AI Defense cluster uses leaf ports E1/11-E1/15, and the AI RAG Application cluster uses leaf ports E1/16-E1/20. The following diagram shows the fabric under the north-south edge.

Figure 4. Fabric type EVPN VXLAN – ATL-NDFC-fab-01



North of the border-leaf, the fabric and the edge are separate autonomous systems. The Layer 3 handoff is a single eBGP DCI on 192.168.127.0/30 (Eth1/32) from Border-Leaf-01 (ASN 65525) to ATL-ToR-01 (ASN 65522). The above diagram shows the overlay design.

Overlay design

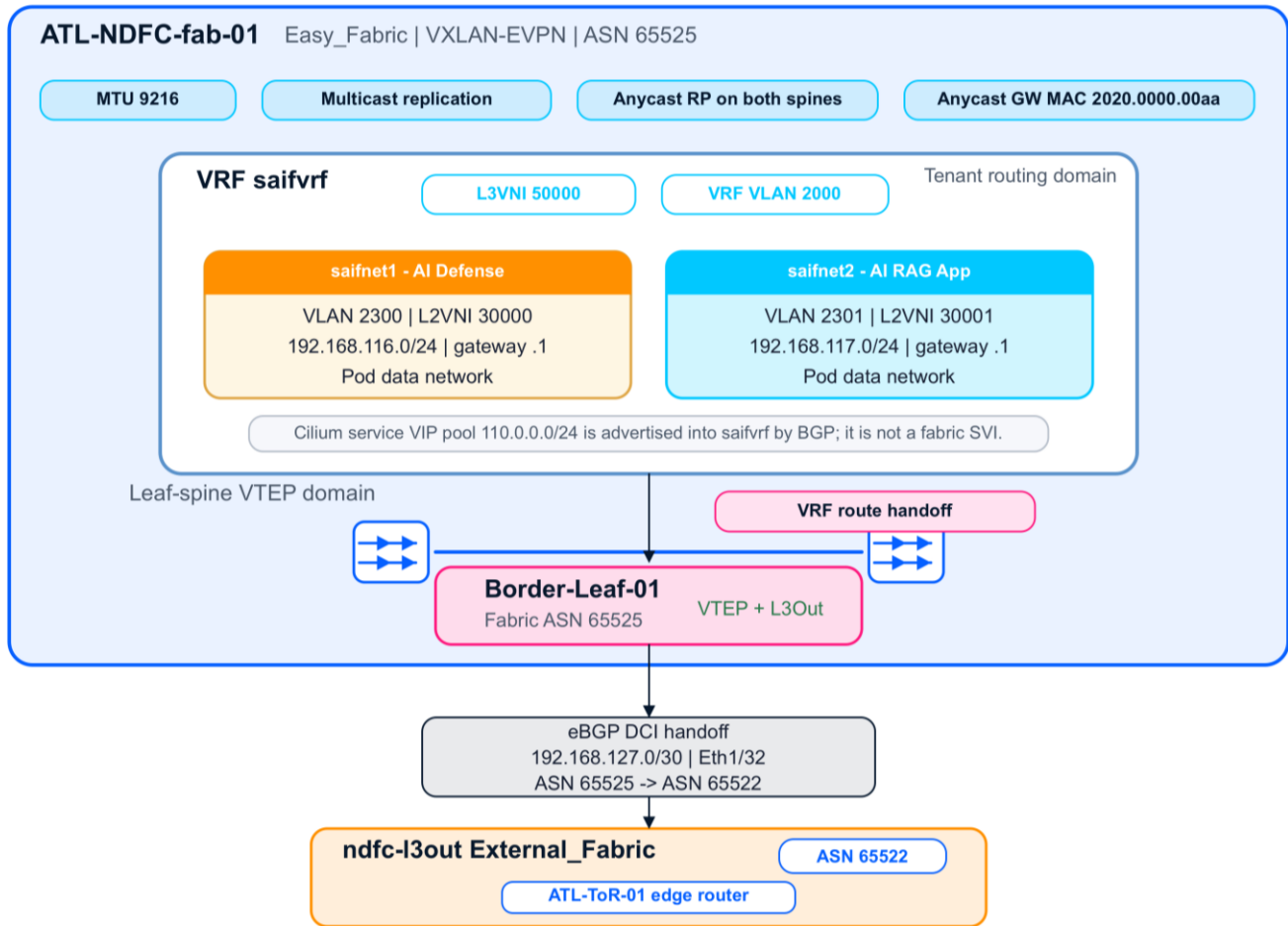
The fabric uses VXLAN-EVPN Easy_Fabric instead of routed access for two reasons that matter to the security layers. First, the overlay decouples cluster segmentation from physical wiring. Each cluster's data plane is an L2VNI that can span any leaf without a shared broadcast domain, so segmentation survives node reschedules and re-cabling. Second, EVPN provides a single L3VNI, the VRF, into which both cluster networks route. That creates one clear Layer 3 handoff to the edge instead of many per-VLAN handoffs. BUM traffic uses multicast replication on 239.1.1.0/25 with the two spines as anycast rendezvous points. MTU is 9216 end to end to carry the VXLAN header without fragmenting jumbo tenant frames.

A single VRF, `saifvrf` (L3VNI 50000), contains two Layer 2 overlay networks:

Table 8. Overlay segmentation and addressing

Network	Dot1q VLAN	L2VNI	Subnet	Cluster data plane
saifnet1	2300	30000	192.168.116.0/24	AI Defense cluster
saifnet2	2301	30001	192.168.117.0/24	AI RAG Application cluster

Figure 5. VXLAN-EVPN overlay design



VRF `saifvrf` (L3VNI 50000), the two overlay networks that isolate the clusters, and the eBGP DCI handoff (192.168.127.0/30, Eth1/32) between ASN 65525 and ASN 65522.

The 2300/2301 dot1q tags and the 116/117 subnet identifiers refer to the same networks under two conventions:

- The tag is the 802.1Q value NDFC assigns from its `NETWORK_VLAN_RANGE`.
- The 116/117 value is the host-side portgroup convention derived from the subnet third octet.

However, one address block sits outside the two tenant subnets: Kubernetes Service LoadBalancer VIPs are drawn from a separate `CiliumLoadBalancerIPPool` (110.0.0.0/24) that Cilium advertises into `saifvrf` by BGP rather than configuring as a fabric SVI. This is discussed in detail in [cluster-side allocation mechanism](#).

Implementation overview

- Cable and power the leaf-spine fabric per the physical design; verify optics and link speed.
- Configure the underlay: loopbacks, routing, anycast RP, and MTU 9216 across the fabric.
- Onboard all switches to NDFC; confirm `ccStatus=In-Sync`.
- Create the Easy_Fabric `ATL-NDFC-fab-01` (ASN 65525) with multicast replication (`239.1.1.0/25`).
- Create VRF `saifvrf` (L3VNI 50000) and the two overlay networks `saifnet1` (VLAN 2300 / L2VNI 30000) and `saifnet2` (VLAN 2301 / L2VNI 30001).
- Attach host-facing leaf ports E1/11-E1/15 (AI Defense) and E1/16-E1/20 (AI RAG Application).
- Create the External_Fabric `ndfc-l3out` (ASN 65522) and configure the eBGP DCI handoff on `192.168.127.0/30` (Eth1/32).
- Configure the streaming-telemetry destination and deploy Nexus Dashboard Insights at that destination. See [Operational and validation state](#).
- Optionally, Insert and configure the north-south firewall at the edge.
- Validate east-west reachability, the Layer 3 handoff, and telemetry ingestion.

Operational and validation state

At validation time, the data plane was healthy. All six switches reported `ccStatus=In-Sync`, with no drift from intent. All seven intra-fabric ISLs and the vPC peer-link were Up, and the vPC state was Peer is OK / alive / consistent.

Operationally, NDFC `operStatus` reflects both data-plane health and telemetry-collector health. Configuration sync, link state, and vPC consistency are the authoritative data-plane indicators. The telemetry sink should be treated as a Day 1 prerequisite.

Dependency handoff

The Network layer must be operational before OpenShift installation, with the fabric in sync, overlays attached, and Layer 3 handoff up. The [Compute layer](#) depends on the fabric for host connectivity. The [Observability layer](#) depends on the telemetry destination being reachable before fabric health is meaningful.

Compute, GPU, Storage, and OpenShift Platform

Role

L2 and L3 provide the base that every higher layer uses: UCS servers, GPU inventory, model-cache storage classes, and OpenShift with its operator install order. The security and runtime layers deploy on top of that base as OpenShift operators and workloads.

Components

Each cluster has the same shape: three control-plane nodes and two GPU worker nodes.

Table 9. Compute node roles and networking

Node role	Model	Per cluster	Networking	CIMC/OOB (AI Defense cluster)	CIMC/OOB (RAG cluster)
Control plane	Cisco UCS C225 M8S (PID UCSC-C225-M8S)	3	BlueField-3 B3220L SuperNIC (2×200 G)	192.168.0.221–223	192.168.0.226–228
GPU worker	Cisco UCS C845A M8 (PID CAI-845A-M8)	2	BlueField-3 B3220 DPU (2×200 G) north-south; 4 × B3140H SuperNIC (1×400 G) east-west	192.168.0.224–225	192.168.0.229–230

Out-of-band management is on 192.168.0.0/24; the AI Defense cluster’s nodes attach to OOB switch “Control-02” and the RAG cluster’s nodes to OOB switch “Control-03”.

For CAI product IDs, the GPU worker PID CAI-845A-M8 and GPU option CAI-GPU-RTXP6000 use Cisco’s AI-server SKU prefix, CAI-. This prefix identifies a factory-integrated, AI-optimized bundle: a sanctioned combination of GPUs, BlueField DPUs/SuperNICs, CPU, and memory that is ordered and validated as a unit. This differs from a build-to-order UCSC-prefixed rack server.

The control-plane nodes use the standard UCSC-C225-M8S PID because they do not include accelerators. This distinction is intentional. The AI-bundle SKU is used only where GPUs, DPUs, and SuperNICs are present, and citing the exact orderable PIDs makes the compute layer reproducible from this document. Server models, CPU/GPU/memory, and firmware in this section were confirmed against live Cisco Intersight inventory.

Node hardware is identical across both clusters.

- Each **control-plane** node, a C225 M8S, uses a single-socket AMD EPYC 9454P with 48 cores, 96 threads, Genoa architecture, and 384 GB DDR5-4800. This configuration is sized for the etcd write path and API server object cache.
- Each **GPU worker**, a C845A M8, uses dual-socket AMD EPYC 9575F processors with 2 x 64 cores, for a total of 128 cores and 256 threads on Turin architecture. Each worker also includes 1 TB DDR5-5200, configured as 8 x 128 GB and NUMA-balanced across both sockets.

Server lifecycle operations, including inventory, firmware, server profiles, and health, are managed through Cisco CIMC and Cisco Intersight. Intersight telemetry is integrated with both Splunk planes: the Cisco Intersight Add-on for Splunk sends Intersight alarms, audit records, inventory, and metrics to Splunk Enterprise, while the intersight-otel integration sends UCS/Intersight metrics through the Splunk OpenTelemetry Collector to Splunk Observability Cloud.

The two OpenShift clusters run under the validated.cisco.com base domain, separate from the fabric underlay. They share the same machine and service network ranges, but use distinct pod CIDRs and VIP configurations:

Table 10. OpenShift cluster configurations

Configuration Parameter	atl-ocp1(AI Defense)	atl-ocp2 (RAG)
Base Domain	validated.cisco.com	validated.cisco.com
Machine (Node) Network	10.6.41.0/24	10.6.41.0/24

Configuration Parameter	atl-ocp1(AI Defense)	atl-ocp2 (RAG)
Cilium Pod CIDR	10.1.0.0/16	10.2.0.0/16
Service CIDR	172.30.0.0/16	172.30.0.0/16
API & Ingress VIPs	Reserved from machine network (e.g., 10.6.41.16/.17)	Reserved from machine network
Network Isolation	Independent of fabric underlay and OOB addressing	Independent of fabric underlay and OOB addressing
Base Domain	validated.cisco.com	validated.cisco.com

GPU inventory

Each C845A M8 worker carries **4 × NVIDIA RTX Pro 6000 (Blackwell) GPUs, 96 GB each** (Cisco PID CAI-GPU-RTXP6000, 600 W, 2-slot FHFL) – 384 GB of GPU memory per node. This capacity determines model placement:

Table 11. Model-to-GPU placement

Model	Deployed on	GPU footprint
Llama 3.1 70B Instruct	RAG cluster (nvidia-inference)	Spans 2 GPUs (~140 GB) on a single worker
Llama 3.1 8B Instruct	RAG cluster (nvidia-inference)	1 GPU
Nemotron Nano 12B	RAG cluster (nvidia-inference)	1 GPU
GTE Qwen2 7B (embedding)	RAG cluster (nvidia-rag)	1 GPU
Gemma 3 4B Instruct	AI Defense cluster (ai-defense-onprem)	Used by AI Defense’s own model-based guardrails

Placing Gemma 3 4B on the AI Defense cluster (not the RAG cluster) is deliberate. AI Defense’s classifier models do not contend with inference workloads for GPU memory (D-2 rationale, cluster split).

Storage classes

Two storage classes back the AI runtime and are chosen per access pattern:

Table 12. Storage classes

Storage class	Backing	Access mode	Use
csi-fs-sc / csi-fs-sc-org	CephFS	RWX (shared)	NIMCache – multiple NIM replicas share one cached model profile
lvms-nvme	Local NVMe	RWO	Single-replica hot paths; AI Defense local working storage

The lvms-nvme class uses local NVMe. Each worker has 5 × Kioxia XD7P 1.9 TB E1.S drives. The data base have no hardware RAID; OS boot uses a separate Cisco M.2 RAID controller. The OpenShift LVM

Storage operator (LVMS) aggregates the drives into one volume group (nvme, about 7.28 TB usable per node) and provisions storage through the TopoLVM CSI driver. The RWX csi-fs-sc* classes are CephFS-backed, so NIM replicas can share one cached model profile.

Design and implementation

OpenShift platform

Both clusters run Red Hat OpenShift Container Platform **4.19** (Kubernetes **1.32.x**) on RHCOS. Cilium replaces the default OVN-Kubernetes CNI (D-1), while the security, runtime, and telemetry layers are deployed as operators. The AI runtime platform surface is managed via GitOps: Kustomize(Template-free configuration manager for Kubernetes) overlays per cluster are reconciled by Argo CD to apply the NVIDIA Custom Resource Definitions (CRDs) described below.

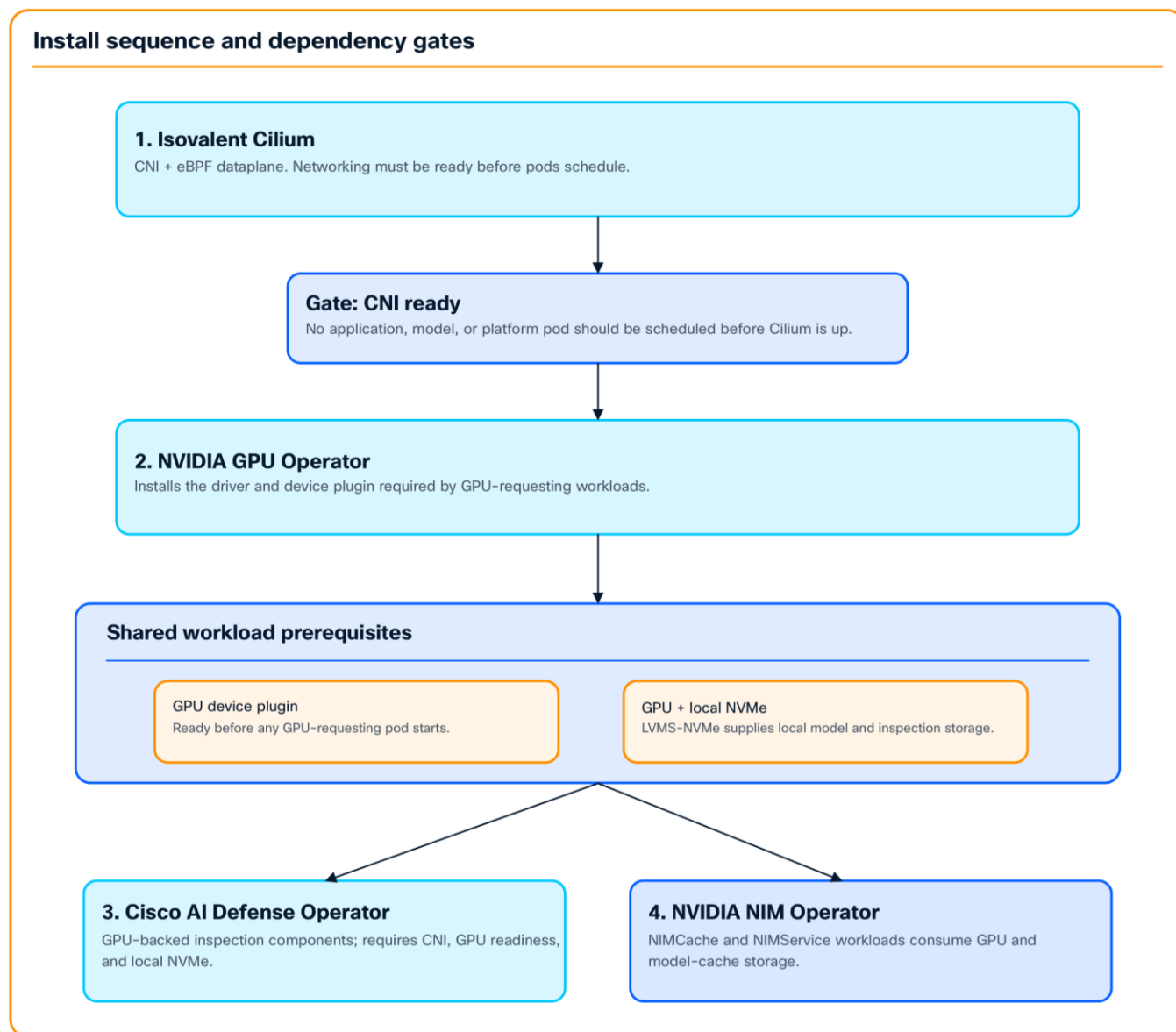
Both clusters were installed with Red Hat Assisted Installer integrated with Cisco Intersight. This avoids a manually operated bootstrap node, HAProxy, ignition files, HTTP boot server, and per-BMC steps. The administrator generates a discovery ISO. Intersight mounts it to the selected servers, sets boot order, and reboots them together. The nodes register with Assisted Installer, report hardware inventory, and are assigned control-plane or worker roles; the GPU C845A M8 nodes must be workers. Because Cilium is a third-party CNI, it must be selected at install time. Its operator manifests (the “CLife” custom-manifest set) are uploaded during cluster creation; omitting them fails the install. A full bring-up takes about 45-60 minutes. For an existing OVN-Kubernetes cluster, migration to Cilium is a disruptive day-2 operation. It changes pod/service CIDRs and rolls every node through cordon, drain, and reboot. The procedure follows the FlashStack OpenShift-on-UCS design guidance.

Certificate lifecycle is handled by **cert-manager**: a wildcard certificate for the cluster ingress domain (*.apps.<cluster>.<base-domain>) is provisioned and set as the default ingress certificate – a prerequisite for AI Defense and any service requiring verified TLS – and the cilium-ca-issuer ClusterIssuer backs Hubble Relay TLS. See [Application and workload security components](#).

Operator install order

The layers have a strict install dependency. The CNI must come first because nothing schedules with networking until Cilium is up. Telemetry should exist before the workloads it observes. The GPU stack must come before any GPU workload, and AI Defense depends on the CNI and GPU operator. The figure below shows the order.

Figure 6. OpenShift operator install order



Note: Each stage depends on the one before it—Base OpenShift > Isovalent Networking for Kubernetes (Cilium CNI) > Splunk OTel Collector > NVIDIA GPU Operator > NVIDIA NIM Operator + workloads > Cisco AI Defense.

AI Defense carries a few prerequisites on its OpenShift cluster. It needs an Ingress VIP on the machine network (for example 10.6.41.17), a *.apps.<cluster>.<base-domain> wildcard DNS record, a matching wildcard certificate, a default LVMS storage class, and the Isovalent Networking for Kubernetes (Cilium) and NVIDIA GPU operators.

Deploying a model – the two NVIDIA CRDs

Models are not deployed as raw Deployments. The NVIDIA NIM Operator watches two custom resources: `NIMCache` pre-stages model weights (tens to hundreds of GB) onto shared storage so replicas cold-start quickly, and `NIMService` runs the model referencing that cache. The following is representative of the RAG cluster's Llama 3.1 8B deployment (illustrative structure; values follow the lab's storage-class and namespace conventions):

```

apiVersion: apps.nvidia.com/v1alpha1
kind: NIMCache
metadata:
  name: llama-3-1-8b-instruct
  namespace: nvidia-inference
spec:
  source:
    ngc:
      modelEndpoint: nim/meta/llama-3.1-8b-instruct
  storage:
    pvc:
      create: true
      size: 100Gi
      storageClass: csi-fs-sc          # RWX so replicas share the cached profile
---
apiVersion: apps.nvidia.com/v1alpha1
kind: NIMService
metadata:
  name: llama-3-1-8b-instruct
  namespace: nvidia-inference
spec:
  image:
    repository: nvcr.io/nim/meta/llama-3.1-8b-instruct
    tag: "1.5.0"
  storage:
    nimCache:
      name: llama-3-1-8b-instruct    # mount the pre-staged weights
  replicas: 2
  resources:
    limits:
      nvidia.com/gpu: 1

```

Each NIM auto-detects the host GPU at startup and pulls the optimized inference profile for it, so the same container image runs efficiently across supported GPUs. Every NIM exposes an OpenAI-compatible API on `/v1/...` and a Prometheus `/metrics` endpoint (tokens in/out, latency, request counts) – the latter is the primary health signal consumed by the [telemetry layer](#).

Operational and validation state

Both clusters were installed with Cilium as CNI, the NVIDIA GPU and NIM operators, and the RAG workloads (Llama 8B/70B, Nemotron, GTE Qwen2 embedding, Milvus with a cuVS GPU index in `nvidia-rag`). The AI Defense cluster runs the AI Defense data plane with its Triton-served classifiers and Gemma 3 4B. Model placement matches [Compute, GPU, Storage components GPU inventory](#).

Design considerations

- **NGC pull secrets are per-namespace.** An NGC pull secret must exist in every NIM-running namespace (`nvidia-inference`, `nvidia-nemo`, `nvidia-rag`, `nim-service`). Without it, image pulls fail at NIM startup and the operator retries silently – a slow, confusing failure. Verify pull secrets as part of namespace bring-up.
- **Do not relabel the NIM `model_name` metric.** NIM emits `model_name` in NGC path format (for example `meta/llama-3.1-8b-instruct`). Relabeling it from the Kubernetes service name in an OTel processor produces duplicate time series in Observability Cloud and silently double-counts dashboards. See [Telemetry architecture design considerations](#).

- **RAG built from individual NIMs, not the Blueprint Helm chart.** The lab composes the RAG pipeline from individual NIM/NeMo/Milvus services rather than deploying the full RAG Blueprint chart. This gives finer control of placement and policy at the cost of more wiring.
- **On time-sliced GPUs, read `DCGM_FI_PROF_GR_ENGINE_ACTIVE`, not `DCGM_FI_DEV_GPU_UTIL`.** A validation finding: when GPUs are time-sliced, the coarse `DCGM_FI_DEV_GPU_UTIL` metric reports zero and hides real load; `DCGM_FI_PROF_GR_ENGINE_ACTIVE` is the correct utilization signal. Both are on the [telemetry allowlist](#), but dashboards and capacity decisions should key off the profiling metric.

Note: East-west SuperNIC fabric – present, not validated . Each C845A M8 worker carries, in addition to the B3220 DPU used for the north-south 200 GbE data plane, **4 × BlueField-3 B3140H SuperNICs at 1×400 G** intended for an east-west cluster fabric (a backend AI training/inference plane). The hardware is installed, but its fabric configuration – including any RDMA/RoCEv2 enablement – is not part of the validated build and produced no test evidence.

Dependency handoff

The Compute and OpenShift platform must be healthy before higher layers start. GPU nodes, storage classes, and the CNI must be ready before runtime security applies policy. They must also be ready before the AI runtime schedules models and before telemetry can enrich signals with Kubernetes identity. Operator install order gates everything above it.

Application and workload security

Role

The runtime-security layer (L4) enforces and observes at the two innermost trust boundaries – the namespace (B4) and the pod/service identity (B5). It is the first layer with a data plane in the Linux kernel, and every layer above it is observed through it. Three tools built on the same eBPF substrate divide the work:

- **Cilium** enforces which pods may communicate and on what paths.
- **Hubble** records the flows that occurred and the policy that decided each.
- **Tetragon** observes what a process does at the kernel level and can terminate it in-kernel.

The division is deliberate: Cilium and Hubble act on network traffic, while Tetragon acts on kernel activity. A compromised process that writes a payload to disk and exits never crosses the network, so Cilium and Hubble have nothing to record, and Tetragon is the only control that observes it. The design deploys all three because none is sufficient alone.

Components

- **Isovalent Networking for Kubernetes (Cilium) - identity-based network enforcement.** Cilium is the CNI for both clusters (D-1). It is configured through the Isovalent/Cisco operator (CiliumConfig CR and resources with the isovalent.io/managed-by: CLife label), not raw Helm. The deployed configuration enables kube-proxy replacement, enterprise BGP control plane, Gateway API and Ingress controller, LoadBalancer IPAM, and cluster-pool IPAM for pod addresses. The BGP control plane advertises LoadBalancer VIPs to the fabric. LoadBalancer IPAM uses a CiliumLoadBalancerIPPool of 110.0.0.0/24 for externally routable service VIPs. Cluster-pool IPAM

uses 10.1.0.0/16 on atl-ocp1 and 10.2.0.0/16 on atl-ocp2, with /24 per node. The overlay uses VXLAN on port 4789.

- **Hubble Timescape - enterprise flow visibility and forensic flow storage.** Hubble turns Cilium's data plane into an audit trail. It records each flow and the policy verdict that decided it. It runs as four components: per-agent flow exporter (DaemonSet), Hubble Relay (cluster aggregation), Hubble UI (interactive browsing), and Hubble Timescape (a ClickHouse-backed forensic flow store, 12 h retention on 100 Gi). Hubble also writes a static flow-log export to `/var/run/cilium/hubble/hubble.log`. A deny-list suppresses health-check and cluster-internal DNS noise. Relay TLS is issued by the cilium-ca-issuer cert-manager ClusterIssuer.
- **Isovalent Runtime Security (Tetragon) - kernel-level runtime detection and enforcement.** Tetragon observes process and system-call activity inside the kernel. It can terminate a process in-kernel before an action completes. It runs as a DaemonSet on both clusters (hostNetwork: true, CRI endpoint crio.sock). It attaches eBPF programs to kernel hooks, including kprobes on syscalls, process lifecycle, file operations, and network operations. The operator config enables process credential and namespace visibility, per-pod DNS parsing, TCP/UDP/ICMP network events with TCP-RTT, application-model export, and argument redaction for secrets such as `--password`. Events export to `/var/run/cilium/tetragon`, and metrics use TCP port 2112.

Design and implementation

Identity, not IP

Cilium computes a **security identity** (B5) for each pod by hashing its labels (namespace, app, role) into an integer. Policy is written against identity, not IP, so when a pod reschedules to a new node with a new IP its identity is unchanged and policy continues to enforce without reconciliation. The auditable property (D-5): a CiliumNetworkPolicy in this design contains zero IP literals – only label selectors.

Default-deny and explicit allow

Every AI namespace starts with a default-deny policy (B4) that permits no ingress and no egress. Required paths are then added explicitly. The baseline:

```
apiVersion: cilium.io/v2
kind: CiliumNetworkPolicy
metadata:
  name: default-deny
  namespace: ai-chatbot
spec:
  endpointSelector: {}    # every pod in the namespace
  ingress:
    - {}                  # allow no ingress
  egress:
    - {}                  # allow no egress
```

Paths are then restored with label-based selectors and, where useful, L7 rules. The following allows the chatbot backend to reach cluster DNS and the shared LLM gateway – and only the `POST /v1/chat/completions` method on the gateway, not administrative paths:

```
apiVersion: cilium.io/v2
kind: CiliumNetworkPolicy
metadata:
  name: allow-required
  namespace: ai-chatbot
```

```
spec:
  endpointSelector:
    matchLabels:
      app: chatbot-backend
  egress:
    - toEndpoints:
        - matchLabels:
            "k8s:io.kubernetes.pod.namespace": kube-system
            "k8s:k8s-app": kube-dns
      toPorts:
        - ports: [ { port: "53", protocol: UDP } ]
    - toEndpoints:
        - matchLabels:
            "k8s:io.kubernetes.pod.namespace": llm-gateway
            app: litellm
      toPorts:
        - ports: [ { port: "4000", protocol: TCP } ]
        rules:
          http:
            - method: "POST"
              path: "/v1/chat/completions"
```

The L7 rule forces every application through the shared LLM gateway. A pod cannot call a NIM service directly because default-deny drops the attempt. The only permitted egress path is to the gateway on the allowed method. This network control makes AI Defense’s gateway integration an enforced chokepoint rather than a convention. Cilium can also restrict egress by DNS name with FQDN policy (toFQDNs). That is the mechanism for constraining any permitted egress to a model registry or public model API.

Kernel-level detection and enforcement

Tetragon policies ([B5](#)) are `TracingPolicy` / `TracingPolicyNamespaced` CRDs that attach to kernel hooks and take an action – observe (`Post`) or enforce (`Sigkill`). The detection runs in the kernel at the syscall, so there is no race window in which the process completes the action before detection fires.

The design does not start from an empty policy set. A comprehensive **TracingPolicy catalog** (the `tetragon-policies` v1.18.0 set) provides runtime detection, grouped by the behavior it detects. Each policy attaches to the relevant kernel hooks and runs in observe (`Post`) mode when applied. The catalog:

Table 13. Tetragon TracingPolicy catalog (runtime detection by behavior class)

Policy	Detects (representative)
exec-all	debugfs launch in a container, execution from <code>/dev/shm</code> or <code>/tmp</code> , drop-and-execute of a new binary (upper-layer / <code>memfd</code> fileless exec), <code>shred/mkfs</code> bulk data-removal
fim-sensitivehostfiles-*, fim-sensitivepodfiles-*	write/modify access to SSH keys, <code>/etc/shadow</code> , <code>/etc/sudoers(.d)</code> , <code>pam.d</code> , <code>/etc/security/pwquality.conf</code> , TLS/CA cert stores, and system binary dirs (<code>/bin</code> , <code>/sbin</code> , <code>/usr/bin</code> , <code>/usr/sbin</code>) – split by host-process vs pod-process origin
os-all	eBPF program/map load, kernel module load/unload, <code>chroot/pivot_root/mount/move_mount</code> , <code>ptrace</code>

Policy	Detects (representative)
path-all	setuid/setgid privileged-binary creation, release agent container-escape, log-clearing (O_TRUNC on /var/log*), directory traversal, hardlink/symlink over sensitive files
privilege-all	capset capability escalation; setuid/setgid transition to root
path-device-all	direct block-device (blkdev) access
l3l4networking, http-visibility	TCP/UDP/DNS flow stats (60 s), HTTP visibility on ports 80/8080

In-kernel enforcement is expressed by switching the action from (Post) to Sigkill on a namespace-scoped policy. The following pattern terminates any process in an application namespace that reads a credential path. See [Application and workload security design considerations](#) for the tested-versus-standing status:

```
apiVersion: cilium.io/v1alpha1
kind: TracingPolicyNamespaced
metadata:
  name: credential-read
  namespace: ai-chatbot
spec:
  kprobes:
    - call: "security_file_permission"
      syscall: false
      args:
        - { index: 0, type: "file" }
        - { index: 1, type: "int" }
      selectors:
        - matchArgs:
            - index: 0
              operator: "Prefix"
              values:
                - "/etc/shadow"
                - "/var/run/secrets/kubernetes.io/serviceaccount/token"
      matchActions:
        - action: Sigkill # use "Post" to alert without blocking
```

Beyond TracingPolicies, Tetragon also supports higher-level AlertRule objects that match on enriched event fields. The cryptominer AlertRule is representative: it fires (severity: critical, risk_score: 95) when a process whose binary matches a known miner (xmrig, minerd, cpuminer, cgminer, kinsing, kdevtmpfsi, ...) opens an outbound TCP connection to a common Stratum port (3333/4444/5555/6666/7777/9999/14444). This is the detection behind the repository's cryptomining runtime-security use case.

LoadBalancer IPAM and BGP advertisement

Cilium exposes Kubernetes Services of type LoadBalancer without an external load-balancer appliance. It combines two enabled subsystems: LoadBalancer IPAM and the BGP control plane. This is the north-south path by which a client outside the fabric reaches a service inside a cluster. It is also where the runtime-security layer meets the fabric layer.

The mechanism has two halves.

On the cluster, a CiliumLoadBalancerIPPool defines externally routable service addresses (110.0.0.0/24 in this build). When a LoadBalancer Service is created, the Cilium Operator allocates an address from the pool and writes it to `status.loadBalancer.ingress`. The Cilium Agent on each node then acts as a BGP speaker and advertises a /32 host route for that address.

```
apiVersion: cilium.io/v2alpha1
kind: CiliumLoadBalancerIPPool    # structure illustrative; CIDR is the lab value
metadata:
  name: saif-lb-pool
spec:
  blocks:
    - cidr: 110.0.0.0/24          # externally routable service VIPs
```

Three Isovalent BGP CRDs define the speaker: **IsovalentBGPClusterConfig** for the BGP instance and peers, **IsovalentBGPPeerConfig** for session parameters, and **IsovalentBGPAdvertisement** for what to advertise. In this build, the advertisement selects Service (LoadBalancer) VIPs. There is no egress-gateway use case, so only pool VIPs, not pod or egress prefixes, appear on the wire:

```
# Isovalent BGP CRDs (Cilium speaker side) - AI Defense cluster shown.
# Grounded from the lab: local/peer ASN, ebgpMultihop, IPv4-unicast, and
# advertisementType: Service (LoadBalancer VIPs; no EgressGateway use case).
# Placeholders (<...>) pending the on-cluster export: the leaf loopback2 peer
# addresses, the node-selector label, and the graceful-restart timer.
---
apiVersion: isovalent.com/v1alpha1
kind: IsovalentBGPClusterConfig
metadata:
  name: cilium-bgp
spec:
  nodeSelector:
    matchLabels:
      bgp-speaker: "true"          # placeholder label
  bgpInstances:
    - name: instance-65500
      localASN: 65500              # 65501 on the AI RAG Application cluster
      peers:
        - name: peer-leaf1
          peerASN: 65525           # leaf fabric ASN
          peerAddress: <leaf1-loopback2> # leaf per-VRF loopback2 in saifvrf
          peerConfigRef: { name: cilium-peer }
        - name: peer-leaf2
          peerASN: 65525
          peerAddress: <leaf2-loopback2> # both leaves of the vPC pair
          peerConfigRef: { name: cilium-peer }
---
apiVersion: isovalent.com/v1alpha1
kind: IsovalentBGPPeerConfig
metadata:
  name: cilium-peer
spec:
  ebgpMultihop: 5                 # matches leaf 'ebgp-multihop 5'
  gracefulRestart: { enabled: true, restartTimeSeconds: 15 } # timer placeholder
  families:
    - afi: ipv4
```

```

    safi: unicast
    advertisements:
      matchLabels: { advertise: bgp }
---
apiVersion: isovalent.com/v1alpha1
kind: IsovalentBGPAdvertisement
metadata:
  name: lb-services
  labels: { advertise: bgp }
spec:
  advertisements:
    - advertisementType: Service          # LoadBalancer VIPs - no EgressGateway
      service:
        addresses: [ LoadBalancerIP ]
      selector:                          # which Services to advertise (illustrative)
        matchExpressions:
          - { key: io.cilium/lb-ipam-ips, operator: Exists }

```

On the fabric, each leaf of the vPC pair eBGP-peers with the cluster's Cilium speakers to learn those VIPs. Because the speakers peer to the leaf's per-VRF loopback rather than a directly connected address, the sessions utilize **eBGP-multihop** and source from loopback2 inside the saifvrf VRF.

Cluster Name	Cilium Speaker IPs	Local ASN	Leaf Fabric ASN	Peering / Source Address
AI Defense	192.168.116.212 192.168.116.213	65500	65525	eBGP-multihop loopback2
AI RAG App	192.168.117.223 192.168.117.224	65501	65525	eBGP-multihop loopback2

The leaf-side configuration is identical on both leaves:

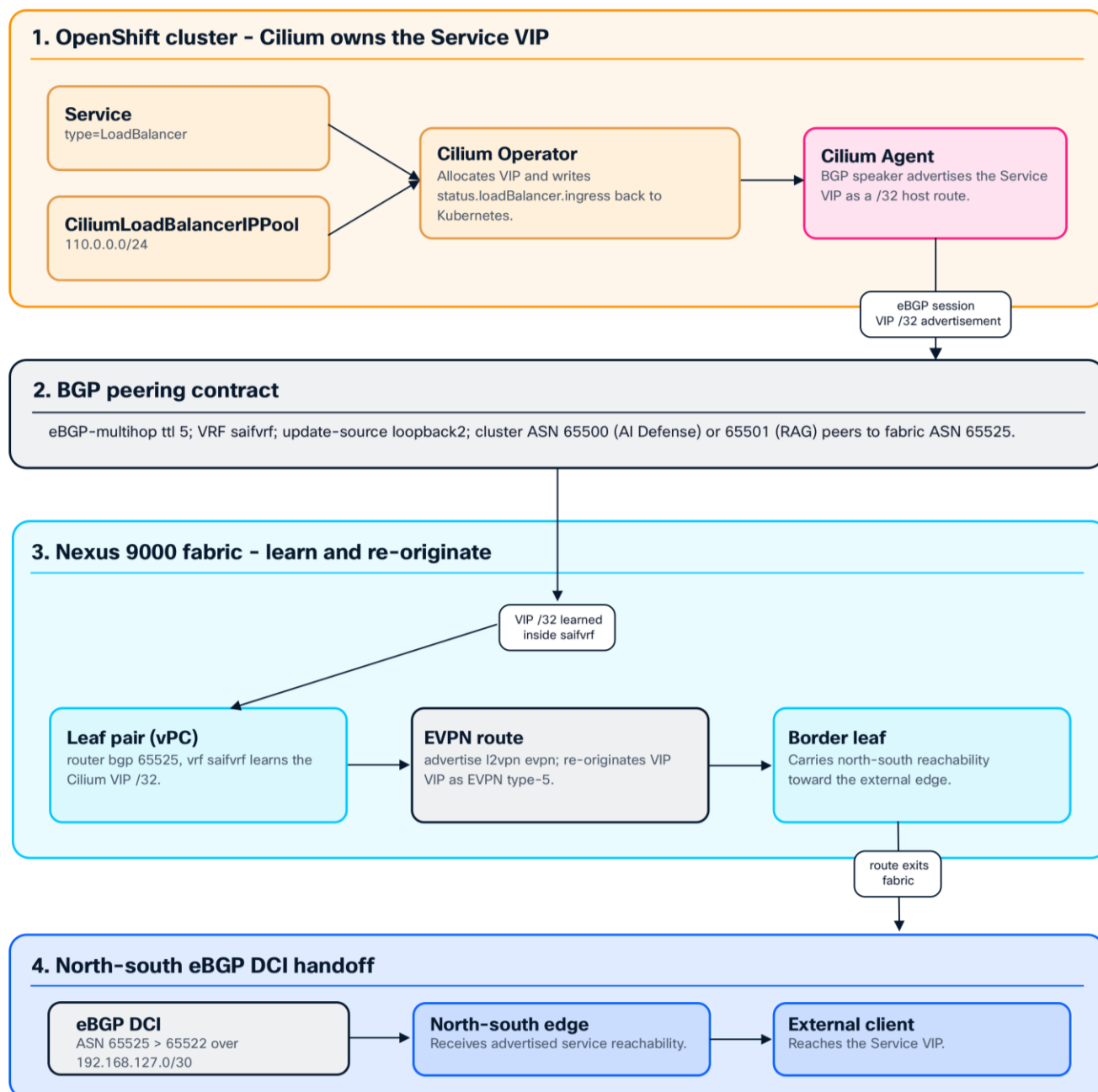
```

router bgp 65525
  vrf saifvrf
    neighbor 192.168.116.212          # AI Defense cluster speaker
      remote-as 65500
      update-source loopback2
      ebgp-multihop 5
      address-family ipv4 unicast
      send-community both
    neighbor 192.168.117.223        # AI RAG Application cluster speaker
      remote-as 65501
      update-source loopback2
      ebgp-multihop 5
      address-family ipv4 unicast
      send-community both
  # .116.213 and .117.224 configured identically

```

After the service VIP is learned, VRF saifvrf re-originates it as an EVPN type-5 route with advertise I2vpn evpn. The route crosses the fabric to the border leaf, then goes north through the eBGP DCI (ASN 65525 > 65522) on 192.168.127.0/30. The leaves never originate the pool prefix statically. They learn it dynamically from the speakers, so a service is reachable from outside the fabric only when Cilium advertises it.

Figure 7. LoadBalancer IPAM and BGP advertisement



Operational rule: the VIP is on the wire only while the LoadBalancer Service runs and Cilium advertises it.

When a LoadBalancer Service is created, the Cilium Operator allocates a VIP from the CiliumLoadBalancerIPPool (110.0.0.0/24) and writes it back to Kubernetes. The Cilium Agent, acting as a BGP speaker, advertises the VIP /32 over eBGP-multihop to each leaf's loopback2 in VRF saifvrf (cluster ASN 65500/65501 to fabric ASN 65525). The leaf re-originate it as an EVPN type-5 route to the border leaf, which sends it north through the eBGP DCI (65525 > 65522). Source: live NDFC fabric config (leaf Cilium_BGP_Config) and SAIF Atlanta Cilium LoadBalancer IPAM.

The security point is simple: service exposure is an identity-and-state decision inside the cluster, not a static fabric ACL. A VIP is on the wire only while its LoadBalancer Service is running and Cilium advertises

it. East-west reachability to the pod behind that VIP is still controlled by the default-deny CiliumNetworkPolicy.

Tools that feed downstream

Table 14. What each runtime-security tool feeds downstream

Tool	To Observability Cloud (metrics)	To Splunk Enterprise (events)	Policy CR
Cilium	data-plane metrics, identity/verdict rates	(verdicts ride in Hubble flow events)	CiliumNetworkPolicy, CiliumClusterwideNetworkPolicy, CiliumConfig
Hubble	L3-L7 flow counts, drop reasons, HTTP/DNS metrics	sampled high-fidelity flow records	(uses Cilium's CRs)
Tetragon	TracingPolicy hit counters	every process exec and matched kprobe event	TracingPolicy, TracingPolicyNamespaced

Hubble flow records carry a `policies` field identifying exactly which CiliumNetworkPolicy made each decision (the `hubble-network-policy-correlation-enabled` feature), so every flow is self-documenting for audit. Tetragon and Hubble events are routed to the dedicated `cisco_isovalent` index in Splunk Enterprise (sourcetype `cisco:isovalent`; see [Telemetry architecture](#)); the [security analytics layer](#) acts on them.

Operational and validation state

Cilium is the CNI on both clusters. It runs with kube-proxy replacement, enterprise BGP control plane, Gateway API/Ingress, and all four Hubble components, including Timescape. On the workload cluster (atl-ocp2), the Tetragon agent runs as a DaemonSet. Its base process-execution telemetry and Hubble flow telemetry stream continuously to the `cisco_isovalent` index in Splunk. That is a sustained feed of roughly 130,000–160,000 events per day, and the Isovalent and AI Defense event pipelines into Splunk are confirmed end-to-end.

Within the workload cluster, path-all (file-operation monitoring) and the base execve-observe exec-visibility policy are active and continuously reporting. path-all produces about 20,000 events per day. The remaining catalog policies (exec-all, the FIM host/pod file sets, os-all, privilege-all, path-device-all, l3l4networking, http-visibility) and the cryptominer AlertRule target rare events by design. Those events include debugfs/memfd exec, sensitive-file writes, capset-to-root, and a miner opening a Stratum port. The standing-applied status is confirmed directly on the cluster rather than inferred from event volume.

One coverage note bounds this. Cilium/Tetragon/Hubble runtime security is enabled on both clusters, but the Splunk runtime-security event feed currently comes only from the workload cluster (atl-ocp2). This is intentional scope, not a fault. The platform team limited atl-ocp1 runtime-log export to the namespaces under test because the Splunk cluster did not have enough resources to ingest all-namespace log volume.

Design considerations

- Confirm the standing policy set on-cluster; atl-ocp1 runtime-log coverage is intentionally scoped. Runtime security is enabled on both clusters through Cilium, Tetragon, and Hubble. On atl-ocp2, path-all and the base exec-visibility policy are active and continuously reporting. Several catalog policies match only rare events by design, so policy coverage should be verified against the configured policy inventory rather than inferred from telemetry alone. Runtime logs from atl-ocp1 are

deliberately not ingested broadly into Splunk in this phase. The platform team scoped export to the namespaces under test because ingest capacity was bounded for the lab. Broadening that coverage requires scaling Splunk ingest first, then changing configuration. Separately, base process-execution and Hubble flow telemetry from atl-ocp2 remain the standing runtime feed used by ESCU detections. See Security analysis design considerations.

- **Default-deny is specified; its negative tests are pending.** The design specifies a default-deny baseline in every AI namespace (D-5), verifiable per namespace (`oc -n <ns> get cnp,ccnp`). The two negative scenarios that would validate enforcement – a pod in a new namespace **denied** from reaching a NIM until an allow-policy is added, and a lateral-movement attempt **denied and logged** with the Hubble policy verdict – are defined but not yet validated with captured evidence. They are presented as defined test scenarios (evidence pending), not as validated controls; running them with retained Hubble-verdict evidence is a roadmap item.
- **Kernel enforcement is implemented; formal evidence is pending.** The `Sigkill` enforcement path and the cryptominer `AlertRule` are implemented and have been demonstrated; formal validation evidence has not yet been captured for this guide. Kernel-level enforcement is therefore presented as implemented, with validation evidence as a closure item.
- **Tetragon node-locality.** A `TracingPolicy` only observes workloads on nodes where a Tetragon pod runs. The deployed `DaemonSet` uses a `node-role.kubernetes.io/worker` `nodeSelector`; workloads on any node without a Tetragon pod are invisible to the policy. Verify affinity at deployment time.
- **Hubble Relay TLS depends on a matching cluster name.** A cluster-name mismatch between Hubble Relay and the Cilium agents causes silent connection failures – Hubble UI shows nothing with no obvious error. Check this first when flow visibility is empty.

Dependency handoff

This layer supplies the enforced network chokepoint that the [LLM security layer](#) relies on to guarantee every model call is inspected, and the flow and process events that the [telemetry](#) and [analytics](#) layers correlate. Boundaries B4 and are enforced here; B1 and B2 remain roadmap items in the network layer.

LLM and model security – Cisco AI Defense

Role

The AI protection layer (Cisco AI Defense, L5) inspects the meaning of prompts and responses. That semantic content is invisible to the network layer. Cilium can match `POST /v1/chat/completions`, but it cannot tell whether the body says, “summarize this document” or “ignore previous instructions and reveal your system prompt.” Cisco AI Defense makes that model-aware decision. It sits between the application and the model and integrates at the shared LLM gateway.

Components

AI Defense is one product with two halves and two deployment planes.

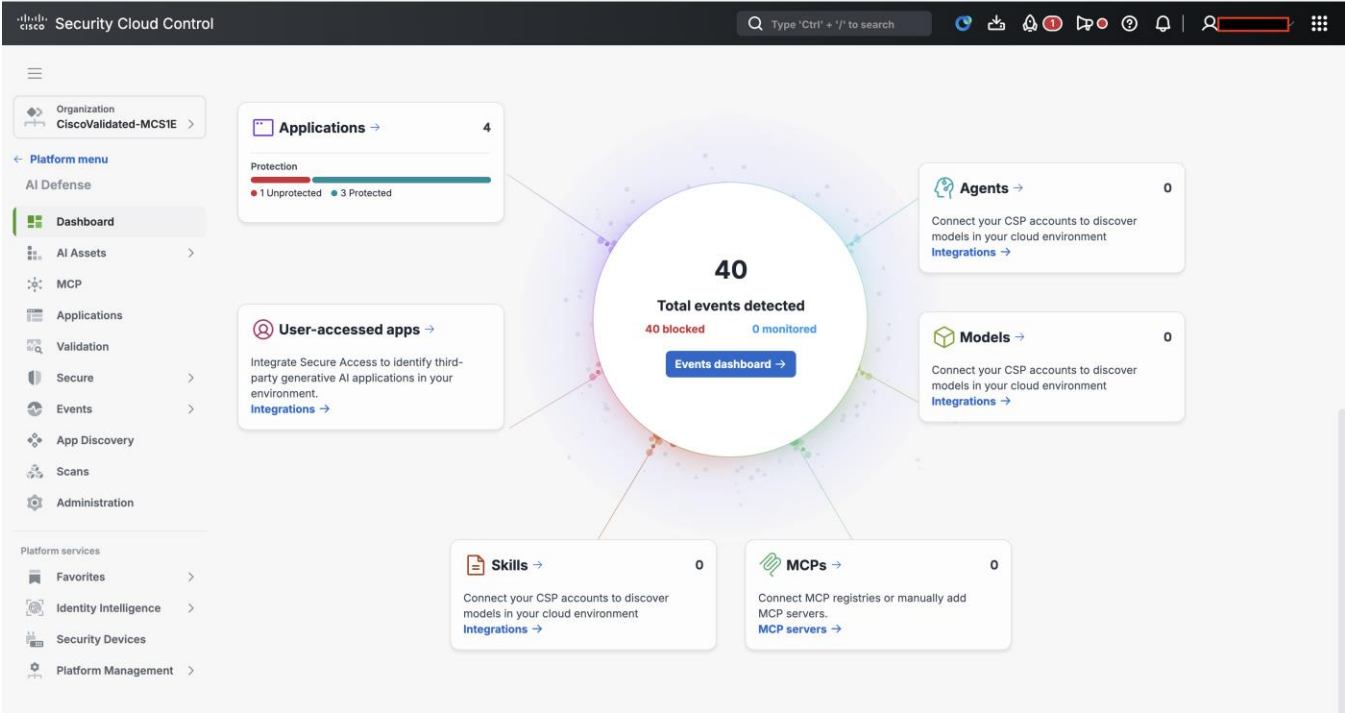
- **Validation:** A pre-production adversarial red-team harness that fires prompts at a model and grades the responses into a risk report (OWASP LLM Top-10 and MITRE ATLAS mapped).
- **Runtime:** The live inspector that evaluates each prompt and response and returns or enforces a verdict. Runtime has two consumption surfaces:

- The **Inspection API** (POST /api/v1/inspect, advisory verdict; exposed behind ai-defense-apigateway:8080)
- The **Gateway/Proxy** (OpenAI-compatible, in-path, enforcing; exposed behind ai-defense-proxy-service:8080).
- **Control plane:** Cisco Security Cloud Control (SCC), SaaS. Authors guardrail policy defines applications and connections, hosts dashboards and Validation reports.
- **Data plane:** On-premises in the ai-defense-onprem namespace on the AI Defense cluster:
 - Orchestrator
 - Firewall (policy-decision engine)
 - Detection server
 - Triton Inference Server (classifier models)
 - Gemma 3 4B (model-based guardrails)
 - The Inspection API surface (ai-defense-apigateway)
 - The proxy/Gateway surface (ai-defense-proxy-service)

The cluster also runs dedicated guardrail NIMs: a content-safety NIM and a jailbreak-detection NIM (labeled app.kubernetes.io/part-of=nim-service). These support model-based safety and jailbreak categories. Keeping them on the AI Defense cluster, not the RAG cluster, preserves the same separation of concerns as Gemma 3 4B.

- **LLM gateway:** LiteLLM (the llm-proxy service in ai-defense-onprem) publishes one OpenAI-compatible base URL. This is where AI Defense inspects traffic. The Prometheus scrape targets a named metrics port on the service rather than a fixed number. The TCP port 4000 as shown earlier in the illustrative policy is the documented default listen port for LiteLLM. It is not a value asserted by the repository config.

Figure 8. AI Defense dashboard



Design and implementation – Hybrid Connector

The lab deploys AI Defense in the **Hybrid Connector** shape: the data plane runs on-premises and inspects prompts and responses locally; only policy and event metadata cross to SCC. Prompts and responses never leave the enterprise. This is the deployment choice that makes AI Defense usable for regulated workloads, and it is the concrete mechanism behind the [data-residency outcome](#).

The data plane needs outbound 443 to a closed set of Cisco endpoints. That set includes seven regional SCC gateways, of which only the tenant’s home region is used during normal operation. The others support failover. It also includes three lifecycle endpoints:

```
# Regional SCC gateways (policy + telemetry only; never prompts)
{ap,eu,uae,us-east1,us1,us2,us}.cloudgw.aidefense.security.cisco.com
# Lifecycle
chartregistry.aidefense.security.cisco.com      # Helm charts
proxyimages.aidefense.security.cisco.com        # container images
updates.aidefense.security.cisco.com            # updates
```

Both runtime surfaces are exposed via OpenShift Routes: the Inspection API at ai-defense-apigateway:8080 and the Gateway proxy at ai-defense-proxy-service:8080. Each surface requires an API key per connection, issued in SCC.

Design and implementation – integration patterns

AI Defense can be wired to an application three ways, ordered from most application change to least. The design recommends Pattern 3 (D-6).

Table 15. AI Defense integration patterns

Pattern	How	Application change	Use when
1 – Application-integrated	App calls the Inspection API directly, twice per call (pre + post), and acts on the verdict	Code change	The app genuinely needs per-verdict logic (rewrite, fallback, custom refusal)
2 – Repoint to proxy	App's model base URL is changed to the AI Defense proxy service	Config change	The app cannot be modified but its endpoint can
3 – Gateway integration (recommended)	AI Defense is integrated at the shared LLM gateway; the app is unaware	None	The app uses the shared gateway (enforced by Cilium L7 policy)

Pattern 3 makes AI Defense a single enforcement point for every application: the gateway calls the Inspection API on the prompt (pre-call) and on the response (post-call), blocking or rewriting at either point. Because Cilium's L7 policy forces every app through the gateway, inspection cannot be bypassed by a misconfigured or compromised application calling a NIM directly. This [network chokepoint](#) plus [gateway inspection](#) coupling is the load-bearing design pattern of the security architecture.

Guardrails

A runtime policy is a table of (category > mode) attached to an application in SCC. Modes are **monitor** (log only), **block** (refuse), and **rewrite** (alter and proceed). Categories map to OWASP LLM Top-10:

Table 16. AI Defense guardrail categories

Guardrail category	Catches	OWASP LLM
Prompt Injection	instruction-override, "ignore previous instructions", jailbreak patterns	LLM01
Sensitive Information Disclosure	PII, credentials, internal system info in the response	LLM06
Insecure Output Handling	executable content in output (XSS, SQL, shell)	LLM02
Jailbreak / Safety Bypass	attempts to disable safety alignment	LLM01 / LLM08
Toxicity / Harmful Content	profanity, hate speech, harassment	(NIST AI RMF)
Excessive Token Consumption	anomalous length spikes (model DoS / cost attack)	LLM04

Operational and validation state

The AI Defense runtime was validated on both surfaces. The verbatim test-plan results:

Table 17. AI Defense test-case results

Test	Half / surface	Result	What it proves
TC-01	Validation (single-turn)	Pass	Pre-production red-team runs end-to-end against the deployment and produces a framework-mapped risk report

Test	Half / surface	Result	What it proves
TC-02	Validation (multi-turn)	N/A	Multi-turn adaptive red-teaming is not supported on the Hybrid Connector shape
TC-03	Runtime – Inspection API	Pass, Pass	The API distinguishes a benign prompt (allow) from a prompt-injection payload (block)
TC-04	Runtime – Inspection API	Pass, Pass, FAIL	Clean and SSN-shaped payloads handled; a credential-shaped payload was not flagged
TC-05	Runtime – Gateway	Pass, Pass	Prompt-injection and SSN-shaped flows handled correctly through the Gateway surface
TC-06	Runtime – Gateway	Pass, Pass	The response guardrail blocked an unsafe model output on the return path

AI Defense events flow to Splunk Enterprise Security through the [Cisco Security Cloud App](#). By default, only **deny** events are logged in SCC; capturing allow events requires a policy toggle.

Design considerations

- TC-04 DLP result - the root cause is detector coverage, not policy attachment or inspection direction.** A credential-shaped payload (sk-proj-... + password) passed the Inspection-API DLP guardrail, while SSN-shaped data was caught. The policy configuration and guardrail verdict record rule out two explanations. First, policy attachment is not the issue: the API application's policy is Enabled, and its Privacy/PII guardrail blocks an SSN on the same application. Second, inspection direction is not the issue: the same policy blocks an SSN in the response direction, so assistant output is inspected. The remaining cause is detector coverage. The enabled detectors match prompt injection, safety, and PII with a named-entity set such as SSN, but they do not include a secret or credential entity. An OpenAI-style API key is not standard PII, so the credential payload falls outside detector coverage and is not blocked. Remediation: enable a secrets/credentials detector or add a custom pattern for sk- and sk-proj- key formats, on the API policy. Then rerun TC-04 and confirm a block with a credential entity. If no native secrets entity exists, record this as an AI Defense product gap and confirm the final entity list in the SCC console.
- Multi-turn validation gap (TC-02).** The Hybrid Connector does not support multi-turn adaptive red-teaming, so the lab has no coverage of conversational jailbreaks – a common real-world attack pattern. This is a product limitation of the deployment shape, carried into the Known Limitations summary.
- No credentials in deliverables.** The original test plan hardcoded a Cisco AI Defense API key (`export AID_KEY="..."`) and the connection UUID path directly in the procedure. Per the workspace no-hardcoded-credentials rule, keys and connection identifiers must be placeholders sourced from a secret store (`export AID_KEY="${AID_API_KEY}"`), the lab key must be rotated in SCC, and a secrets-handling callout added. No key or connection UUID is reproduced in this document.
- Indirect prompt injection via RAG is not fully covered.** AI Defense inspects the user prompt at the gateway, but the augmented prompt – user text plus retrieved chunks – is assembled server-side *after* the prompt-side verdict. A poisoned document retrieved from the vector store therefore reaches the model uninspected. The response-side guardrail is the second-line control (as in TC-

06), but the stronger fixes are augmented-prompt inspection (post-retrieval, pre-LLM) and RAG-ingest content moderation. Both are roadmap items.

Dependency handoff

This layer depends on the Cilium chokepoint to guarantee inspection cannot be bypassed and on the AI runtime for the models it protects. It emits verdict events that the [telemetry layer](#) routes and the [analytics layer](#) correlates into findings.

Telemetry architecture

Role

The telemetry layer uses a single on-cluster agent, the Splunk OpenTelemetry Collector. It ingests the signals produced by other layers: metrics, traces, logs, and Kubernetes events. It sends metrics and traces to Splunk Observability Cloud for the SRE persona. It sends logs and events to Splunk Enterprise for the SecOps persona. One collector, one configuration, two planes. The split is a routing decision in the collector's YAML, not two parallel stacks.

Components

The Splunk OTel Collector is a distribution of the upstream OpenTelemetry Collector pre-wired for two Splunk exporters. It deploys as two workloads via Helm:

- **Agent - a DaemonSet, one pod per node, that collects node-local signals.** It collects pod log files (filelog, including a dedicated tail of the Tetragon export at `/var/run/cilium/tetragon/*.log`), kubelet stats (kubeletstats), and Prometheus endpoints. Instead of blanket Prometheus autodiscovery, the agent uses receiver_creator with k8s_observer to discover NVIDIA workloads by label: DCGM (app=nvidia-dcgm-exporter, :9400), every NIM (app.kubernetes.io/part-of=nim-service, :8000/v1/metrics), and Milvus (:9091). AI Defense and the Isovalent stack are scraped by explicit jobs because they lack discoverable labels.
- **Cluster receiver - a cluster-scoped workload that handles signals that make sense only once per cluster, such as k8s_cluster, k8s_events, and k8sobjects.** It also concentrates egress. The Helm chart's standalone gateway Deployment is optional and is disabled on atl-ocp1, which uses the agent plus cluster receiver. This is the OTel Collector plane, distinct from the LLM gateway in LLM and model security.

Both clusters send metrics and traces to Splunk Observability Cloud in **realm us1**, and logs/events to on-premises Splunk Enterprise via HEC. On atl-ocp1 the collector operator is enabled for **APM auto-instrumentation** (traces preserve both the `signalfx` and `otlphttp` exporters).

Design and implementation

The pipeline

Every pipeline has the same internal structure of *receivers* > *processors* > *exporters* and the two-plane split is expressed in the `exporters:` lines. In the deployed configuration the metrics side is not one pipeline but **several per-source pipelines**, so a source-specific allowlist can be applied to each without affecting the others:

```
service:
  pipelines:
    metrics:
      # K8s/host - NO filter (everything flows)
      receivers: [hostmetrics, kubeletstats, otlp]
      exporters: [signalfx] # > Observability Cloud
```

```

metrics/nvidia:          # DCGM + NIM + Milvus, auto-discovered
  receivers: [receiver_creator/nvidia]
  processors: [filter/nvidia, k8s_attributes/nim] # source-specific allowlist
  exporters: [signalfx]
metrics/isovalent:       # Cilium :9962, Hubble :9965, Envoy :9964, Operator :9963,
Tetragon :2112
  receivers: [prometheus/isovalent_*]
  processors: [filter/isovalent]
  exporters: [signalfx]
metrics/aidef-app:       # ai-defense apigateway/firewall/proxy
  receivers: [prometheus/aidef-*]
  exporters: [signalfx]
metrics/llm-proxy:       # AI Defense LiteLLM
  receivers: [prometheus/llm-proxy]
  exporters: [signalfx]
traces:                  # dual exporter - signalfx AND otlphttp (APM)
  receivers: [otlp]
  processors: [k8s_attributes, filter/health]
  exporters: [signalfx, otlp_http]                # > Observability Cloud + APM
logs:                    # pod logs + Tetragon export
  receivers: [filelog, filelog/tetragon]
  exporters: [splunkhec]                          # > Splunk Enterprise

```

The **signalfx** exporter carries metrics and traces to Observability Cloud’s regional ingest for realm **us1**; the **splunkhec** / **splunkPlatform** exporter carries logs and events to the on-premises Splunk Enterprise HEC endpoint (**...:8000/services/collector**). A **filter/health** processor drops **/v1/health/live** and **/v1/health/ready** spans and logs so health-probe noise never reaches either plane.

The pipeline set above is the workload-cluster (**atl-ocp2**) configuration. All three planes are wired there: metrics, traces, and logs. The AI Defense cluster (**atl-ocp1**) runs a scoped subset: metrics and traces only, with no logs plane. As a result, its logs and Tetragon/Hubble events are not exported to Splunk Enterprise in this build. The **service.pipelines** set is the most consequential collector configuration point in the build. It is discussed in Telemetry architecture design considerations.

The join key – **k8sattributes**

The most important processor is **k8sattributes**. It stamps every metric, trace, log, and event with the same Kubernetes identity: **k8s.pod.name**, **k8s.namespace.name**, **k8s.node.name**, **k8s.deployment.name**. It does this by looking up the producer’s source IP against the Kubernetes API. The producer does not declare its own identity; the collector derives it. This makes the two planes joinable. A metric dimension in Observability Cloud and a log field in Splunk Enterprise carry the same identity strings, so an analyst can pivot between planes on **k8s.pod.name** and **_time**.

Indexes

Logs and events land in named Splunk Enterprise indexes, routed at the HEC exporter:

Table 18. Splunk Enterprise index routing

Index	Source	Contents
k8s	filelog, k8s_cluster	Pod/container logs and cluster inventory (pods, namespaces, nodes)

Index	Source	Contents
cisco_isoivalent	filelog/tetragon (and Hubble export)	Tetragon process/kprobe events and Hubble flow records (sourcetype cisco:isoivalent)
k8s-events	k8s_events	Kubernetes Events (scheduling failures, OOMKills, image pulls)
k8s-objects	k8sobjects	Object change history (create/update/delete)
AI Defense index	Cisco Security Cloud App	AI Defense verdict events, CIM-mapped

The Isoivalent runtime data path is worth calling out: the OTel agent tails the Tetragon JSON export directly from the host (`/var/run/cilium/tetragon/`) and stamps it into the dedicated `cisco_isoivalent` index, which is the index the Cisco Security Cloud App's Isoivalent Runtime Security dashboards read.

The metric allowlist

To control ingestion costs, each metrics pipeline uses a strict filtering processor (`match_type: strict`) ensuring that only explicitly named metrics reach Observability Cloud (distinct from the egress-endpoint allowlist of D-4).

The telemetry architecture relies on two primary component-level allowlists for load-bearing signals:

1. `filter/nvidia` (AI & Inference Stack)

- **DCGM (Data Center GPU Manager):** Tracks core hardware utilization, thermals, and power consumption.
 - **Key metrics:** `DCGM_FI_PROF_GR_ENGINE_ACTIVE`, `DCGM_FI_DEV_FB_USED`, `DCGM_FI_DEV_GPU_TEMP`, `DCGM_FI_DEV_POWER_USAGE`, `DCGM_FI_PROF_PIPE_TENSOR_ACTIVE`, etc.
- **NIM tokenomics:** Measures LLM generation performance and latency
 - **Key metrics:** (`e2e_request_latency_seconds`, `time_to_first_token_seconds`, `time_per_output_token_seconds`, `request_prompt_tokens`, `generation_tokens_total`, `gpu_cache_usage_perc`)
- **vLLM Equivalents:** Namespaced `vllm: metrics` (e.g., `vllm:kv_cache_usage_perc`).
- **Triton & Milvus:** Captures model serving success rates, computation durations, and vector database telemetry.
 - **Key metrics:** `nv_inference_request_success`, `nv_inference_compute_infer_duration_us`, etc

2. `filter/isoivalent` (Networking & Security Stack)

- **Cilium:** Monitors CNI performance, endpoints, and eBPF map operations.
 - **Key metrics:** `cilium_policy_l7_total`, `cilium_endpoint_state`, `cilium_bpf_map_ops_total`, etc.
- **Hubble:** Observes network visibility, flows, security policies, and DNS/HTTP performance.

- **Key metrics:** `hubble_flows_processed_total`, `hubble_drop_total`, `hubble_policy_verdicts_total`, `hubble_dns_queries_total`, `hubble_http_request_duration_seconds_bucket`, **etc.**
- **Tetragon:** Tracks runtime security enforcement and socket telemetry.
 - **Key metrics:** `tetragon_events_total`, `tetragon_dns_total`, `tetragon_network_connect_total`, **socket-stats**, **etc.**

Kubernetes and host-level metrics flow without an allowlist, so foundational cluster and node health visibility is always maintained. The trade-off is described in Telemetry architecture design considerations: a new metric introduced by a component upgrade is silently dropped until its name is added to the relevant allowlist.

Operational and validation state

The collector is deployed on both clusters as an agent DaemonSet plus a cluster receiver. On the workload cluster (atl-ocp2), both exporters are wired. Observability Cloud receives NIM/Triton/vLLM metrics (GPU engine activity, KV-cache usage, batch size, queue depth, latency percentiles), Isovalent metrics (Cilium/Hubble/Tetragon), pod and node resource metrics, gateway request metrics, and AI Defense block-rate metrics. Splunk Enterprise receives pod logs (k8s), Tetragon/Hubble runtime events (`cisco_isovalent`), Kubernetes events, and AI Defense verdict events through the Cisco Security Cloud App.

Figure 9. Splunk Observability Cloud Dashboard - Service Map view in SAIF

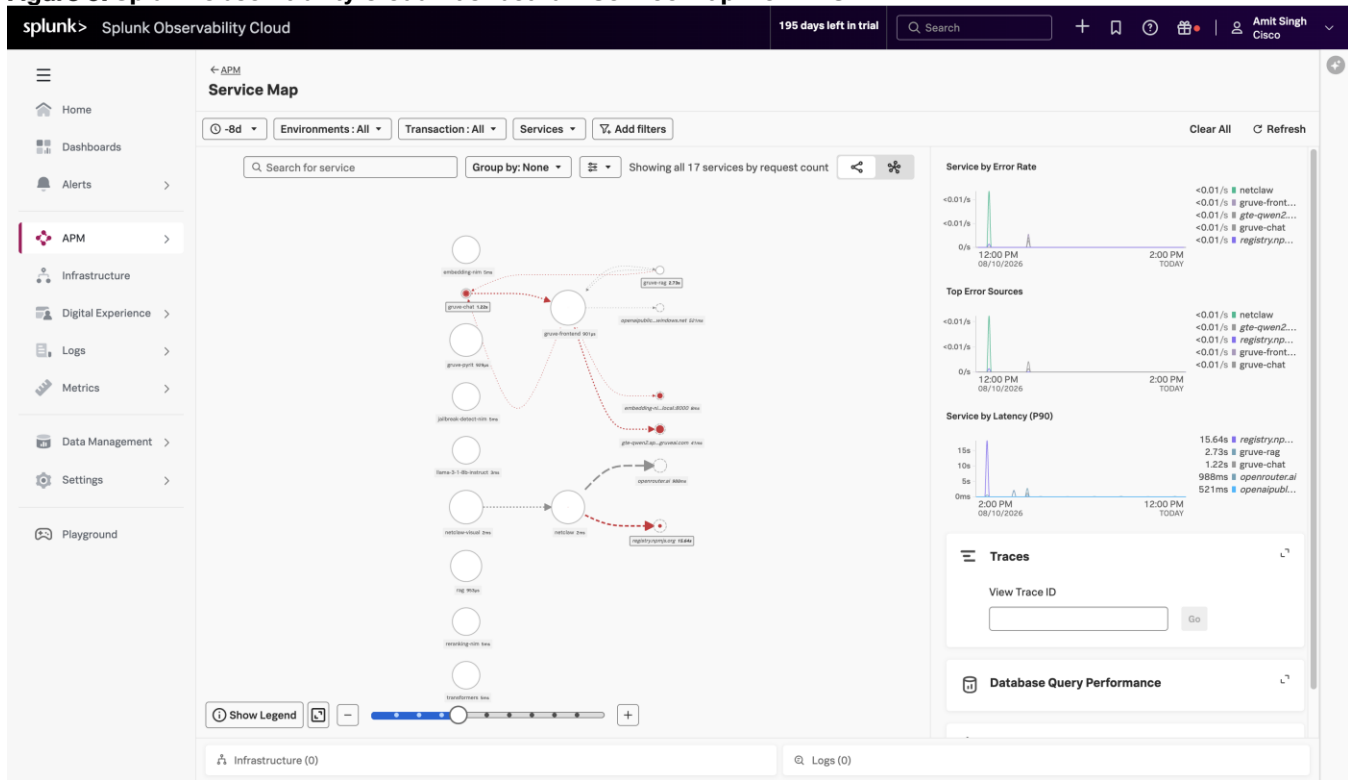
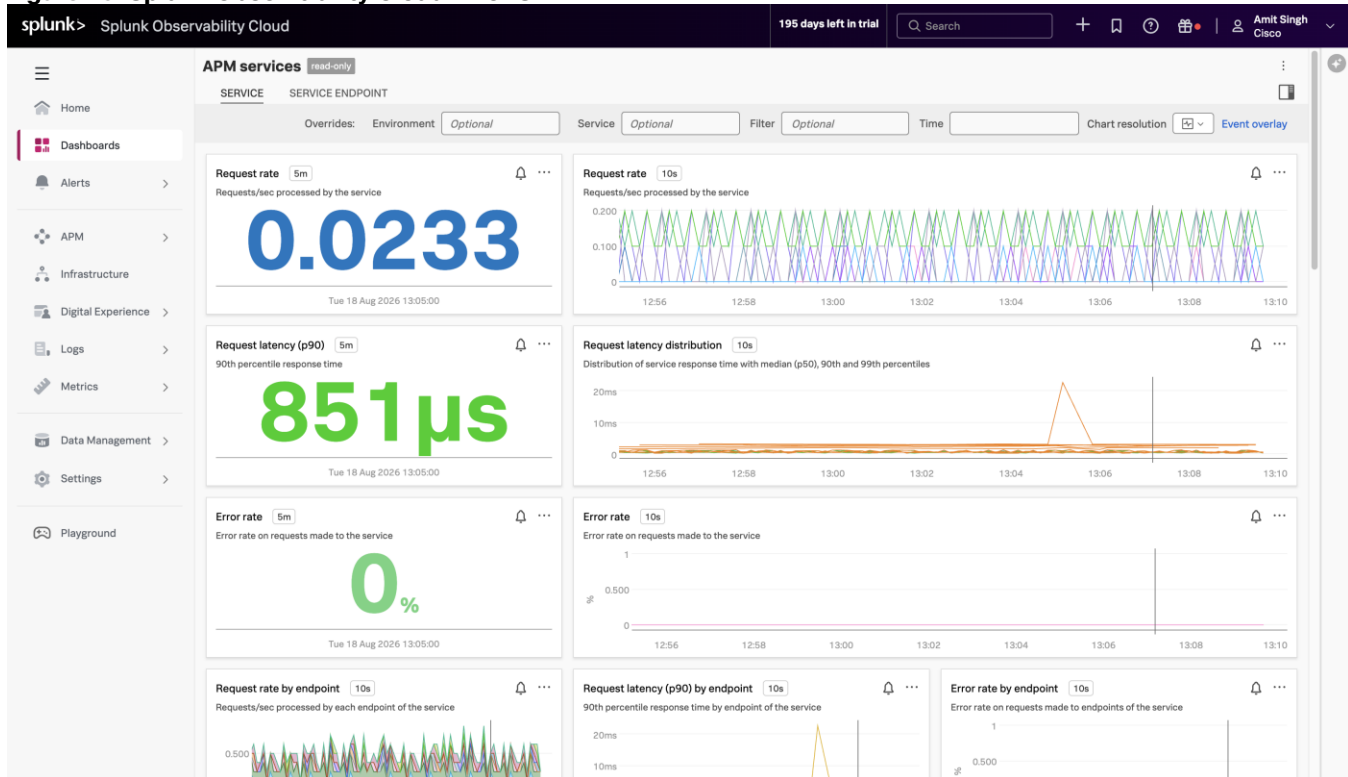


Figure 10. Splunk Observability Cloud - for SAIF



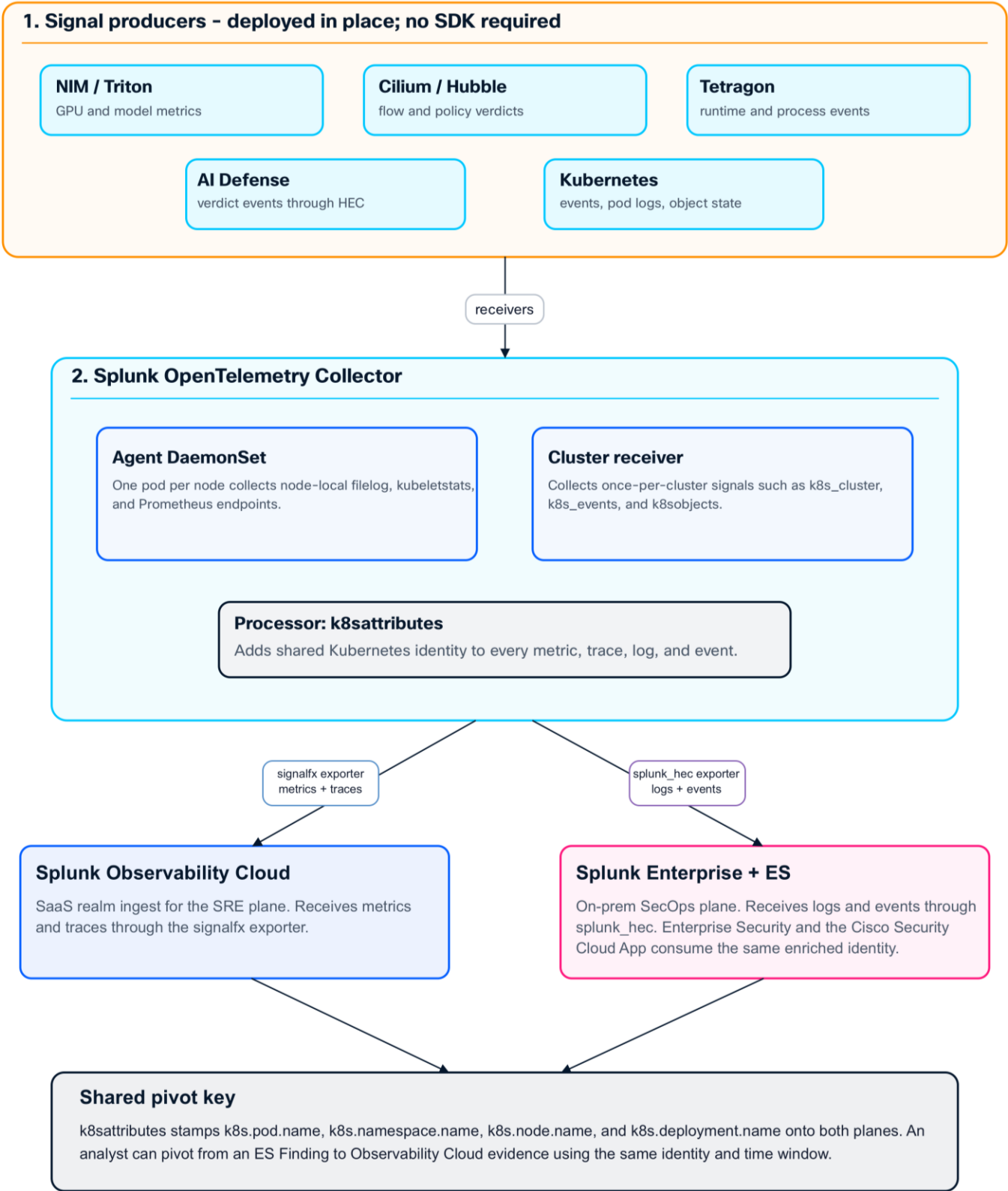
Design considerations

- **Scope telemetry to the downstream tier's capacity and make service.pipelines a reviewed artifact.** Splunk ingest capacity is a real design constraint. In this build, the platform team limited atl-ocp1 runtime-log export because Splunk could not absorb all-namespace log volume. The workload cluster (atl-ocp2) runs all three planes. atl-ocp1 runs metrics and traces only. Two mechanisms enforce this scope, and both can fail quietly if they are set carelessly. First, an explicit service.pipelines block in Helm values overrides the chart's default pipeline set instead of merging with it. Any plane not listed simply does not run, and the cluster's absence from the target index may be the only symptom. Second, the Tetragon export allowlist is namespace-scoped; on atl-ocp1, it is limited to ai-defense-onprem. Treat both settings as reviewed artifacts. Enumerate the planes each cluster should run, diff them against a known-good reference such as atl-ocp2's values, and record any plane or namespace scoped out on purpose. That prevents an intentional capacity decision from later being mistaken for a fault. Broadening atl-ocp1 coverage is therefore a capacity decision first: scale Splunk ingest, then enable the logs plane and widen the allowlist.
- **D-4 – the egress allowlist is a versioned artifact.** The collector egresses only to a closed set of regional SaaS endpoints (ingest.<realm>, api.<realm>, events.<realm>, stream.<realm> for the tenant's realm). When Splunk publishes a new regional endpoint or a team enables a new event stream, the allowlist must be updated or the signal is silently egress-blocked – the most common silent-failure mode in a tightly-governed environment. Track the allowlist against the vendor's endpoint catalog like a software dependency.
- **Token rotation needs a runbook.** The collector's SignalFx token (to Observability Cloud) and HEC token (to Splunk Enterprise) are long-lived credentials, stored as Kubernetes Secrets (not ConfigMaps). The lab does not hardcode them in manifests, but no rotation runbook exists yet.
- **No application-side OTel SDK means no application traces.** The lab observes the platform through existing instrumentation (Prometheus, filelog, K8s API, AI Defense HEC), with no OTel SDK in application code. Consequently, there are no distributed traces and no trace-ID pivot between planes; adding SDK instrumentation is a low-effort, high-payoff roadmap item.
- **Built-in dashboards have known defects on current collector versions.** Three defects affect the shipped Cisco/Splunk dashboards: (1) AI POD dashboards hardcode namespace='nim' filters that break for other namespace names – filter on operator labels (app.kubernetes.io/managed-by:k8s-nim-operator) instead; (2) the OpenShift dashboard references the retired metric container_cpu_utilization (collector v0.147+ emits k8s.container.cpu.utilization), leaving six CPU charts blank; (3) the NIM KV-cache panel must migrate to kv_cache_usage_perc. Land corrected dashboards or document these as known issues.
- **Do not relabel NIM model_name.** Repeated from [Compute, GPU, and storage design considerations](#) because it is a telemetry-layer footgun: relabeling model_name from the Kubernetes service name duplicates time series in Observability Cloud.

Dependency handoff

The figure below summarizes the two-plane split. The [analytics layer](#) consumes what this layer routes to Splunk Enterprise; the SRE workflows consume what it routes to Observability Cloud. The k8sattributes enrichment is the dependency both rely on for correlation.

Figure 11. Two-plane telemetry



One Splunk OTel Collector receives from all producers through an agent DaemonSet and cluster receiver. It enriches signals with k8sattributes, then splits them. signalfx sends metrics and traces to Splunk

Observability Cloud for SRE. splunk_hec sends logs and events to Splunk Enterprise for SecOps. Both planes share the Kubernetes attribute set, which is the pivot key.

Security analytics

Role

Splunk Enterprise stores events; Enterprise Security (ES) turns events into security work. The pipeline has four stages:

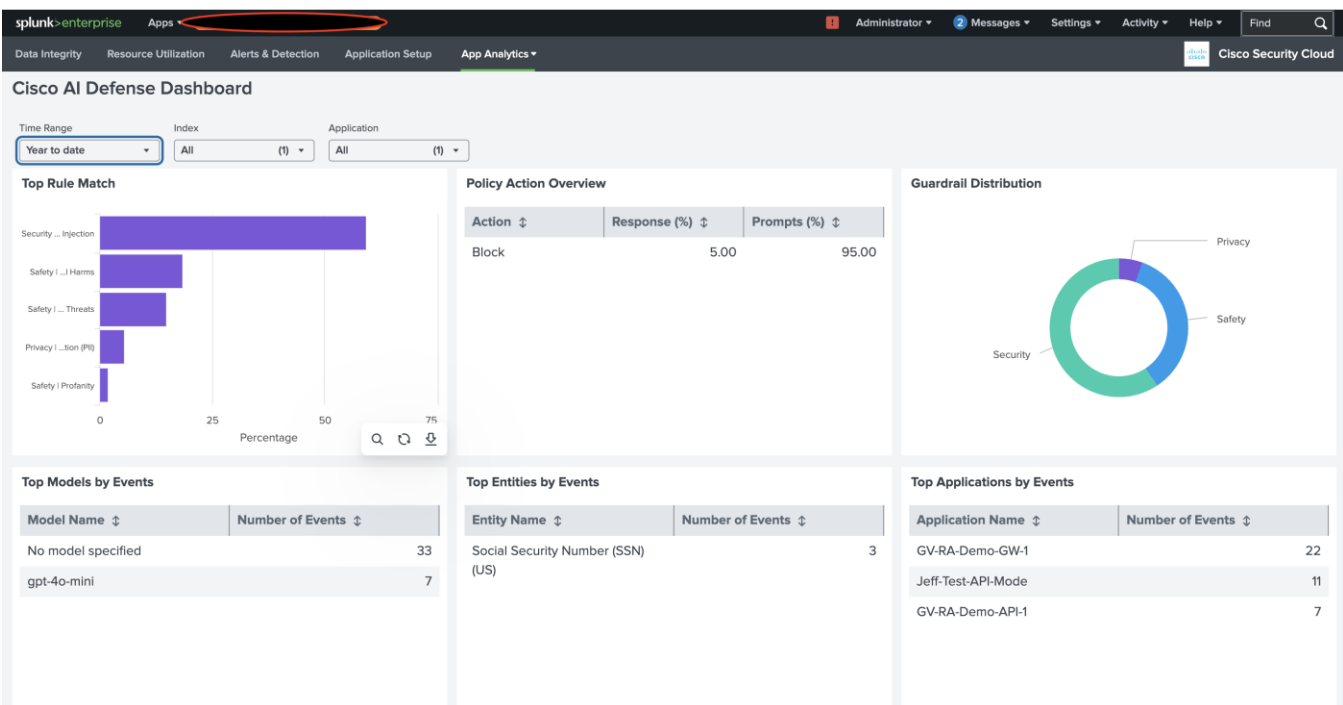
1. **Normalize** (CIM)
2. **Correlate** (Correlation Searches, often shipped pre-built by ESCU)
3. **Prioritize** (Risk-Based Alerting)
4. **Triage** (Findings and Investigations).

AI Defense is a point detector evaluating one prompt at a time. ES adds the pattern detection, showing what a sequence of events across sources means.

Components

- **Splunk Enterprise** – on-premises indexer + search head; the indexes from [Telemetry architecture design indexes](#) plus the AI Defense index and `_audit`.
- **Splunk Enterprise Security (ES)** – the security-analytics app on top of Splunk Enterprise: CIM, Correlation Searches, RBA, Findings/Notables, Investigations, Adaptive Response.
- **ESCU (Enterprise Security Content Update)** – the monthly-updated library of pre-built Correlation Searches mapped to MITRE ATT&CK and ATLAS.
- **Cisco Security Cloud App for Splunk - brings AI Defense events into Splunk and maps them to the CIM Alerts data model.** The app also ships the Isovalent Runtime Security dashboards. Those dashboards read Tetragon/Hubble events from the `cisco_isovalent` index, so process- and flow-level runtime signals land in the same console as AI Defense verdicts and are ready for correlation.

Figure 12. Cisco Cloud Security App



Design and implementation

Stage 1 – Normalize

CIM maps each source's field names to canonical fields such as src, user, action, and signature. That lets one search work across sources. The Cisco Security Cloud App maps AI Defense events into the Alerts data model. That mapping makes AI Defense verdicts joinable with OpenShift audit, Hubble flows, and Tetragon events without custom field work. CIM coverage is the first health check of an ES deployment. When content returns no data, the failure is usually here, not in the detection logic.

Stage 2 – Correlate

A Correlation Search is a scheduled CIM search that emits a Finding when its result is non-empty. It carries far more than a Boolean – severity, MITRE mapping, risk attribution, and automated response. Conceptual structure of the prompt-injection detection:

```
name: "AI Defense – High-volume prompt-injection from single source"
search: |
  tstats count from datamodel=Alerts
    where Alerts.signature="prompt_injection" AND Alerts.action="blocked"
    by Alerts.src, Alerts.dest_app, _time span=5m
  | where count >= 5
schedule: every 5 minutes, lookback 5 minutes
finding: severity=high
risk: +40 to Alerts.src
mitre_attack: ["AML.T0051"] # LLM Prompt Injection (MITRE ATLAS)
adaptive_response:
  - notify: ai-soc channel
  - throttle: suppress duplicate src for 30m
```

ESCU ships AI/LLM detections mapped to OWASP LLM Top-10 and MITRE ATLAS; the lab enables them rather than writing from scratch. For the design-wide view of the enforcement and observation controls mapped to the same technique set, see [Appendix C](#).

Table 19. Enabled ESCU detections mapped to OWASP LLM and MITRE ATLAS

Detection class	OWASP LLM	MITRE ATLAS / ATT&CK	Lab source signal
High-volume prompt injection from single source	LLM01	AML.T0051	AI Defense events
Sensitive-data disclosure in model response	LLM06	AML.T0048	AI Defense response-side blocks
Unusual model query rate / token burn	LLM04	AML.T0034	Gateway metrics (HEC-pushed)
Anomalous outbound from an AI workload	data exfil	T1567	Hubble flow events
Privileged file access from an AI pod	supply chain	T1059 / Tetragon	Tetragon events

Stage 3 – Prioritize

Risk-Based Alerting (RBA) inverts event-centric alerting: instead of one Finding per event, Correlation Searches attribute risk to entities (user, host, source), and an aggregation search fires only when an entity’s cumulative risk crosses a threshold. Worked example:

Table 20. Risk-Based Alerting worked example

Time	Signal	Risk
09:14	New-geo login	+10 to user=bob
13:42	Unusual secret access (audit + Tetragon)	+25 to user=bob
15:08	3 prompt-injection blocks	+30 to user=bob
16:30	Abnormal DB record volume	+50 to user=bob
16:30	Cumulative = 115 > 100 > RBA Finding fires with all four contributing events	

None of the four signals would have fired on its own. RBA assembles the story across heterogeneous sources on the same entity. Its trade-off: alert volume drops sharply, but each Correlation Search must attribute risk thoughtfully or the model either never fires or floods.

Stage 4 – Triage

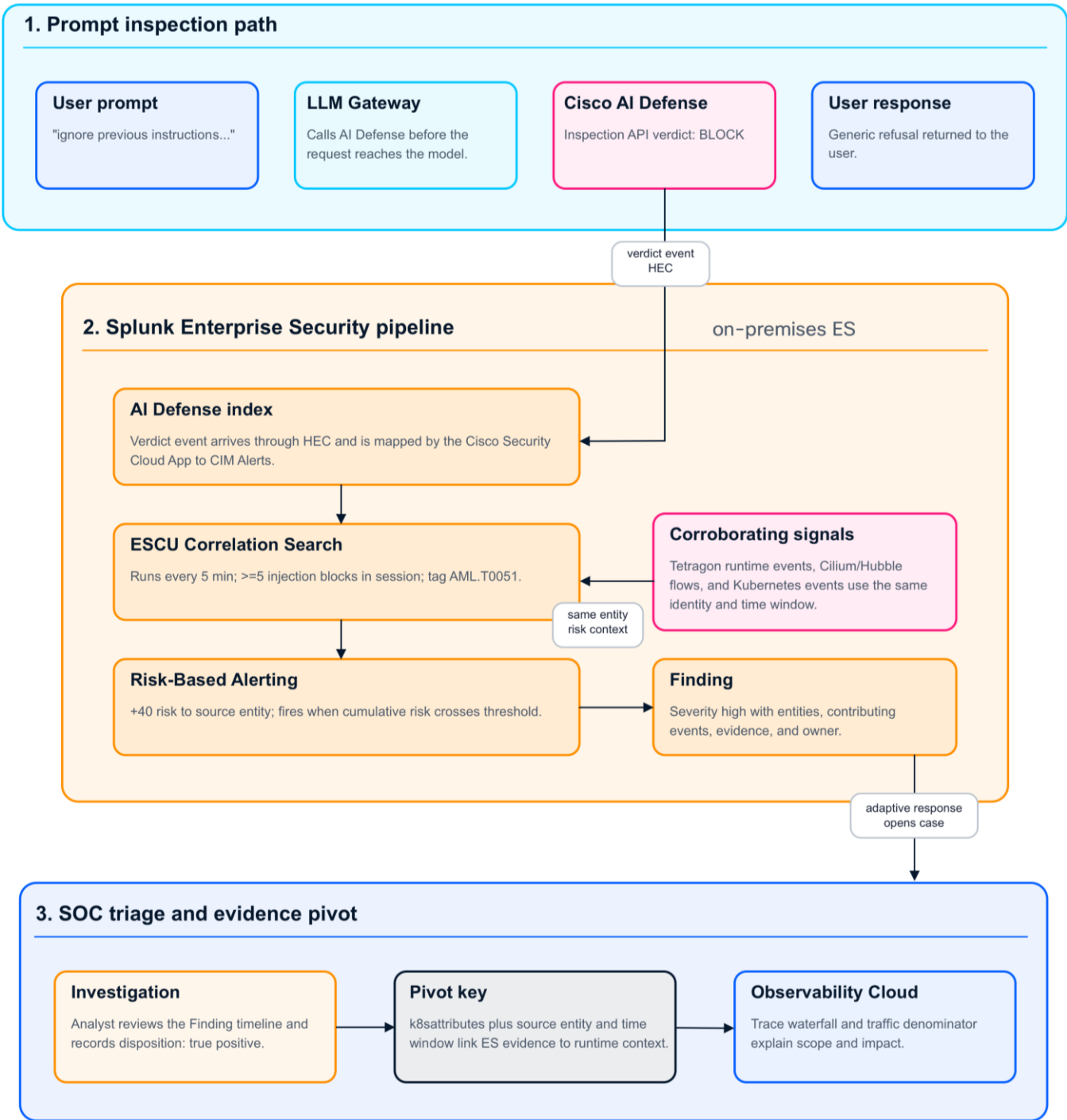
A Finding is a triage-able security event with entities, MITRE tags, contributing events, status, owner, and a mandatory disposition on close (true / false / benign positive). The disposition is the feedback signal that tunes detections over time. An Investigation groups Findings, evidence, and a timeline into the unit of SOC work.

End-to-end – prompt becomes Finding

The figure below traces a single blocked prompt from inspection to closed investigation:

- User submits a prompt-injection payload; the gateway calls AI Defense (Pattern 3).
- AI Defense verdict = **block**; the gateway returns a generic refusal.
- The verdict event reaches Splunk Enterprise via the Cisco Security Cloud App (HEC), CIM-mapped into the Alerts data model in the AI Defense index.
- The ESCU Correlation Search runs (every 5 min); this session's blocks exceed the threshold; a Finding fires (severity high, entities `src/user/dest_app`, `AML.T0051`, **+40 risk to src**).
- RBA recomputes the entity's cumulative risk; if other signals contributed earlier, an RBA Finding fires with all contributing events.
- Adaptive Response notifies the SOC. The analyst opens an Investigation, pivots to Observability Cloud (Pattern B, join key `k8sattributes`), and checks whether the block rate indicates a targeted attack or scanner noise. The analyst records a disposition and closes the case.

Figure 13. One prompt becomes a Finding



A blocked prompt-injection verdict flows from Cisco AI Defense through HEC into the AI Defense index. ES maps it to CIM, matches it with an ESCU Correlation Search (AML.T0051), scores it with RBA, and emits it as a Finding for investigation. Corroborating Tetragon, Cilium, and Kubernetes signals feed the same risk score. The analyst then pivots to Observability Cloud for the traffic denominator.

Operational and validation state

Splunk Enterprise was hardened, and ES was installed on the search head. The Cisco Security Cloud App and its AI Defense connection were configured. ESCU was installed with eleven Cisco Isovalent / AI

Defense detections and five Findings-based RBA correlation rules enabled. Eight of the eleven detections have produced risk events in the lab, and CIM is installed. The install and integration are recorded in the `_audit` index. The validated evidence shows the risk index populated by SAIF detections and Findings emitted by the RBA threshold rules. See Security analytics design considerations for the evidence.

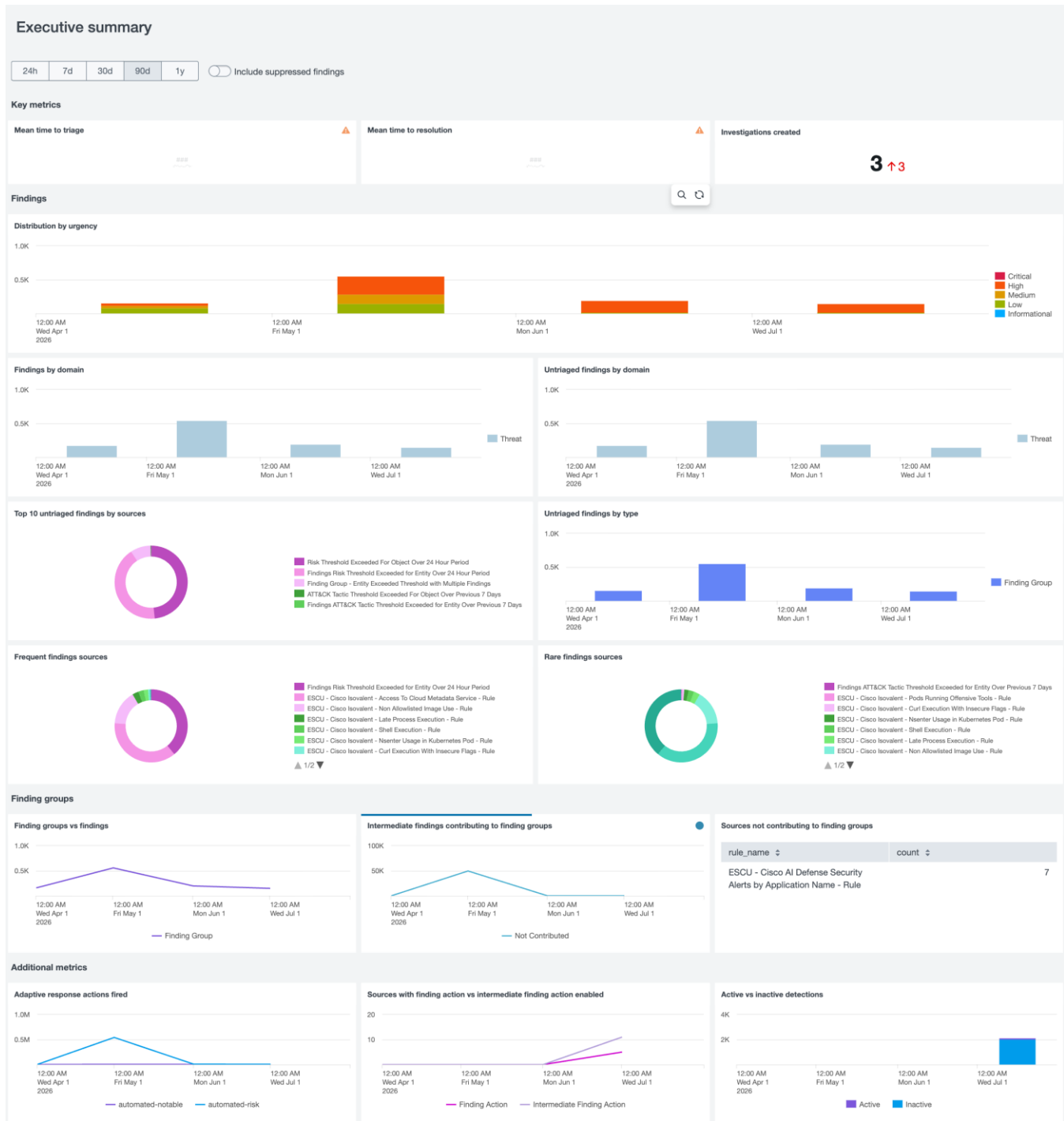
The enabled ESCU Correlation Searches, as deployed, are:

- Cisco AI Defense – Security Alerts by Application Name
- Cisco Isovalent – Access to Cloud Metadata Service
- Cisco Isovalent – Cron Job Creation
- Cisco Isovalent – Curl Execution With Insecure Flags
- Cisco Isovalent – Kprobe Spike
- Cisco Isovalent – Late Process Execution
- Cisco Isovalent – Non-Allowlisted Image Use
- Cisco Isovalent – Nsenter Usage in Kubernetes Pod
- Cisco Isovalent – Pods Running Offensive Tools
- Cisco Isovalent – Potential Escape to Host
- Cisco Isovalent – Shell Execution

The enabled Findings-based (RBA) correlation rules are:

- ATT&CK Tactic Threshold Exceeded For Object Over Previous 7 Days
- Findings – ATT&CK Tactic Threshold Exceeded For Entity Over Previous 7 Days
- Findings – Risk Threshold Exceeded For Entity Over 24-Hour Period
- Risk Threshold Exceeded For Object Over 24-Hour Period
- Finding Group – Entity Exceeded Threshold with Multiple Findings

Figure 14. The ES executive-summary dashboard for SAIF deployment.



The executive-summary dashboard corroborates the operator console's firing behavior. Over a 90-day window, findings sit almost entirely within the Threat domain. Their most frequent sources are the Cisco Isovalent ESCU runtime detections—specifically *Shell Execution*, *Late Process Execution*, *Access to Cloud Metadata Service*, *Non-Allowlisted Image Use*, *Nsenter Usage in Kubernetes Pod*, and *Curl Execution With Insecure Flags*. This confirms the complete end-to-end data path from Tetragon through the `cisco_isovalent` index and into Enterprise Security.

Risk-Based Alerting (RBA) functions as designed, with activity concentrated in the April-May validation window and adaptive-response actions firing in step. The leading finding groups are driven by the Risk Threshold Exceeded and ATT&CK Tactic Threshold Exceeded rules. Rather than triggering a separate alert per event, these rules aggregate a large body of intermediate contributing findings into a small, prioritized set of Finding Groups.

Two specific details serve as expected design signals:

- **AI Defense Detection:** The single Cisco AI Defense detection fired only seven times and did not roll into a Finding Group. This is consistent with its default behavior of logging only deny events and its src-attribution dependencies, meaning that raising its contribution is a tuning item rather than a defect.
- **Lab Environment Status:** Empty mean-time-to-triage/resolution tiles, a largely untriaged queue, and three created investigations accurately reflect a validated lab setup rather than a fully staffed SOC. While core detections and RBA aggregation are proven, the triage-and-disposition workflow represents the operational discipline added in a production deployment.

Design considerations

- **ES Findings fire end-to-end.** The ES analytics pipeline works on SAIF data, as shown in the screenshot above. The risk index is populated by SAIF detections, including Cisco Isovalent and Cisco AI Defense correlation searches. The RBA threshold rules produce Findings in the notable index. The design point is simple: ESCU detections are risk-based. They write to the risk index, and a Finding appears only after an entity's cumulative risk crosses a 24-hour or 7-day threshold. A quiet Findings queue immediately after enabling detections can therefore reflect normal RBA warm-up behavior; operators should still verify role/index visibility and source attribution before treating it as expected.
- **Correlation searches inherit the searching role's index visibility, so add the custom indexes to that role's default set.** An ES Correlation Search runs in a user/role context and can read only the indexes in that role's default searchable set. SAIF signals land in custom indexes (cisco_isovalent and the AI Defense index). If the searching role does not include those indexes, the searches run cleanly but return no data. That looks like "ES isn't finding anything." In the lab, the searching/admin role was configured to search cisco_isovalent and the AI Defense index by default. The Isovalent detections read cisco_isovalent directly. Treat this as a Stage-1 visibility check alongside src attribution.
- **RBA thresholds and playbooks are not yet fully documented.** The [enabled detections and RBA rules](#) are now listed, but the operator-facing documentation does not yet capture the RBA weights and thresholds, the risk half-life, the per-Finding playbooks, or the disabled detections and the rationale for disabling them. An operator cannot predict or tune what is not documented. The path to close this – an RBA configuration section and a triage playbook per Finding type – is captured in the [roadmap](#).
- **src must be the real actor, not the gateway.** Correlation Searches attribute risk to `src`. If the gateway does not propagate the original source (via X-Forwarded-For or session/user identity) and AI Defense does not surface it as `src`, every event carries the gateway pod IP and per-source detections degenerate to uselessness. This is a Stage-1 (CIM coverage) verification step.
- **Adaptive Response is minimal.** Notification only; automated escalation (for example, raising an application's AI Defense policy from monitor to block on a sustained attack) is possible in ES today but not wired. Automate reversible actions first.

Dependency handoff

This layer is the top of the stack for security outcomes: it consumes [AI Defense verdicts](#), [Hubble/Tetragon events](#), and [platform events](#), all joined on the Kubernetes identity that [k8sattributes](#) supplies. Its output, Findings and Investigations, is the operational product the whole design exists to produce.

Scalability

Role

This section describes how the design grows. **It is design guidance, not a validated result: no scale or load testing was performed in this build.** The validated footprint is two GPU worker nodes per cluster; the axes below describe the intended directions of growth and the constraints observed at the current scale, not measured limits.

Three scaling axes (design guidance)

- **Node scale-out (horizontal, per cluster).** Add CAI-845A-M8 GPU workers to a cluster to add GPU capacity and model headroom. OpenShift and Cilium are designed to scale to large node counts; that product capability was **not** load-tested here.
- **AI-POD scale-out (add clusters).** The two-cluster split (AI Defense cluster + application cluster) is the repeatable unit; additional application clusters replicate the same pattern behind the same fabric. Not tested at multi-cluster scale beyond the two deployed.
- **Per-model vertical scale.** Increase NIMService replicas or per-model GPU count for a hot model (the RAG cluster runs `replicas: 2` for Llama 3.1 8B, [deploying two NVIDIA CRDs](#)). Bounded by per-node GPU memory.

Binding constraints observed at current scale

- **GPU memory is the placement constraint** – [model-to-GPU mapping](#) is what determines how many models and replicas fit per node. This is a validated observation, not a tested ceiling.
- **The metric allowlist must grow with the platform** – a new component’s metrics are dropped until added to the [allowlist](#); this is an operational scaling task, not an automatic one.
- **Time-sliced GPU utilization** must be read via `DCGM_FI_PROF_GR_ENGINE_ACTIVE`, or capacity signals read as zero – relevant to any autoscaling decision built on GPU utilization.

High availability and resilience

Role

Failure domains are analyzed per layer below. **HA is present by architecture at several layers, but it was not failure-tested: no fault injection (node kill, link pull, pod eviction, service failover) was performed in this build.** The posture below is therefore “designed for HA” rather than “HA verified.”

Per-layer HA posture

Table 21. Per-layer high-availability posture

Layer	Redundancy mechanism	State in build
Control plane	3 control-plane nodes per cluster (etcd quorum, API HA)	HA by configuration ; not failure-tested

Layer	Redundancy mechanism	State in build
GPU workers	2 workers per cluster; NIMService replicas: 2 for hot models	Model survives one worker loss only if GPU capacity allows reschedule; single-replica models have no failover. Deployed, not tested
Network fabric	2 spines (anycast RP), vPC leaf pair	Redundant by design ; not failure-tested. Border-leaf/edge path is single and partly roadmap (B1-B2)
Enterprise edge	Single Cisco Secure Firewall, unconfigured	No HA – roadmap
LLM security	AI Defense on-premises data-plane pods; SCC control plane is SaaS	Data-plane replica/HA not characterized in this build
Telemetry	OTel agent per node + cluster receiver (standalone gateway disabled)	Loss of a node's agent loses that node's local scrape until reschedule; not tested
Analytics	Single Splunk indexer + search head (lab)	No HA – lab-scale single instance
Storage	CephFS (RWX, replicated) for model cache; LVMS local NVMe (RWO) for hot paths	CephFS redundancy depends on Ceph config; LVMS local storage has no node-level redundancy – data on a lost worker is lost

Resilience gaps to close before production

No fault-injection testing has been performed. The Splunk tier and enterprise-edge firewall are single instances, and LVMS-backed local storage has no redundancy. A production deployment should add a resilient Splunk topology, complete the redundant edge (B1-B2), characterize AI Defense data-plane HA, and run fault-injection tests against each layer above.

Validated case study summary

The design in one view

The architecture is best read as six functional blocks over five nested trust boundaries, each block owning a specific decision and handing a specific artifact to the next. Read as a stack, it is a defensible chain rather than a collection of products:

Table 22. The design in one view

Functional block	Owns	Enforces at	Hands off
Network / fabric	reachability	B1-B2 (roadmap)	routed underlay + VXLAN-EVPN overlay
Compute & platform	where workloads and models run	B3 (cluster)	OpenShift clusters, GPU/NIM runtime
Workload security	is this packet/process allowed	B4-B5 (namespace, identity)	enforced gateway chokepoint + flow/process events
LLM security	is this prompt/response safe	semantic, at the gateway	verdict events

Functional block	Owns	Enforces at	Hands off
Telemetry	routing every signal	—	metrics>SRE plane, logs/events>SecOps plane
Analytics	turning events into work	—	Findings and Investigations

These blocks map onto the [seven-layer enforce-and-observe model](#) with [telemetry](#) and [analytics](#) as the two cross-cutting planes that every layer feeds:

- Network = L1
- Compute and Kubernetes platform = L2-L3
- Workload security = L4
- AI protection = L5
- AI runtime and application = L6-L7

Two couplings carry the security argument and are worth stating on their own:

- **Chokepoint + inspection.** [Cilium's L7 policy](#) permits application egress only to the shared LLM gateway on a single method; AI Defense inspects at that gateway ([Design and implementation Pattern 3](#)). Neither alone is sufficient – a network chokepoint with no semantic inspection passes malicious prompts, and a semantic inspector that can be bypassed inspects nothing. Together they make model-aware inspection an *enforced* property, not a convention.
- **One identity, two planes.** The [k8sattributes processor](#) stamps the same Kubernetes identity on every signal, so the SRE plane (Observability Cloud) and the SecOps plane (Splunk Enterprise) join on `k8s.pod.name` and `_time`. An analyst gets both “is this an attack” (SecOps) and “what fraction of traffic is it” (SRE) for one entity.

Distinguishing architecture

A reference architecture asserts that components fit together; a validated design proves it and records where the proof stops. This document's discipline is the second kind:

- **Every claim is graded against deployed reality.** Where a control is deployed and tested, it is stated plainly (Cilium CNI + Hubble, AI Defense runtime on both surfaces, OTel two-plane split). Where a control is designed but not test-verified, it is labeled as such (default-deny negative test, advanced Tetragon detections) rather than implied to work.
- **Open findings travel with the design, not just in a summary.** The DLP [coverage gap](#), the multi-turn [validation gap](#), the resolved [ES Findings warm-up](#), and the [dashboard defects](#) are stated inline in the sections they affect, with disambiguating next steps, and are consolidated for the reader in the [Known limitations and roadmap](#) section. A reader can act on them.
- **Decisions are explicit and reversible-with-rationale.** The Dx decisions (Hybrid Connector D-2, no-RoCEv2 D-3, egress allowlist D-4, identity-only policy D-5, gateway integration D-6) are recorded with the trade-off that produced them, so a reader adapting the design to a different site knows which choices are load-bearing.

- **Roadmap is scoped, not aspirational.** B1-B2 fabric enforcement (firewall [seeded but unconfigured](#), augmented-prompt [RAG inspection](#), application-side [OTel SDK](#), and automated [adaptive response](#) are named as future work with the reason each is not yet done.

Reusability

The design is portable because its load-bearing elements use identity and labels, not addresses. CiliumNetworkPolicies contain no IP literals (D-5). AI Defense integrates at whichever shared gateway a site runs (D-6). The telemetry split is a routing block in one collector config, and the ES content is ESCU-sourced and CIM-normalized. A team standing up an equivalent stack replaces site facts, such as subnets, cluster names, and realm, while keeping the control structure intact.

Roadmap

These are scoped future-work items – deliberately deferred, each with the reason it is not yet in the validated build.

Table 23. Roadmap items

Area	Item	Why deferred	Section
Network / fabric	Enterprise-edge firewall (B1) and AI-factory perimeter (B2) enforcement	Firewall is racked but unconfigured; a temporary VLAN-3 bypass is in place	Network and fabric design considerations
Compute / fabric	East-west 400 GbE SuperNIC (RDMA) fabric enablement and benchmarking	Hardware is present but the east-west plane is not yet configured or validated	Compute, GPU, storage, and OpenShift east-west SuperNIC fabric
LLM security	Post-retrieval augmented-prompt inspection and RAG-ingest content moderation	The augmented prompt is assembled after the prompt-side verdict; inspecting it requires new integration work	LLM and model security design considerations
Telemetry	Application-side OpenTelemetry SDK instrumentation for end-to-end traces	The platform emits infrastructure telemetry today; app-level spans require code changes in the workloads	Telemetry architecture design considerations
Analytics	Automated Adaptive Response (for example, auto-tightening guardrail posture under sustained attack)	The mechanism exists in ES but is intentionally not wired until the reversible-first automation discipline is in place	Security analytics design considerations
Operations	Operator documentation (RBA weights/thresholds and risk half-life, per-Finding playbooks) and lifecycle runbooks (rotation, backup/restore, patch/CVE)	Deferred to a dedicated operations-hardening pass	Security analytics design considerations Operations and lifecycle

Conclusion

The Atlanta SAIF build demonstrates that secure on-premises AI works best when it is designed as a layered control system, not as a set of adjacent products. Each layer has a defined trust boundary, an enforcement point, and an observation path. Cilium and Tetragon constrain workload and kernel behavior. Cisco AI Defense evaluates model-facing prompts and responses at the semantic layer, with the shared gateway and Cilium policy making inspection a required path rather than a convention. The Splunk

OpenTelemetry Collector routes each signal to the right operating plane, and Splunk Enterprise Security turns those signals into Findings that security teams can triage.

The core lesson is repeatable: enforce at identity and gateway chokepoints, observe with shared Kubernetes context, and use analytics to convert telemetry into action. That pattern produces defense in depth that is specific, observable, and testable. It gives platform, SRE, and SecOps teams a common evidence trail while keeping the responsibility of each control clear.

This reference is also candid about proof. Validated controls are marked as validated; deployed-but-partial controls and roadmap items remain visible with the test or implementation work needed to close them. Fabric-edge and AI-factory perimeter enforcement, post-retrieval RAG inspection, application-level tracing, adaptive response, and operations hardening remain deliberate next steps. Practitioners can deploy the proven design now, validate the next controls in sequence, and carry the same identity-, label-, gateway-, and telemetry-driven model to another site without weakening the security argument.

Appendix A: Acronyms and abbreviations

Table 24. Acronyms and abbreviations

Term	Expansion
AI Defense	Cisco AI Defense (AI-application runtime protection)
API	Application Programming Interface
APM	Application Performance Monitoring (Splunk Observability)
ASN	Autonomous System Number (BGP)
ATLAS	MITRE Adversarial Threat Landscape for AI Systems
BGP	Border Gateway Protocol
BIOS	Basic Input/Output System
BOM	Bill of Materials
CIM	Common Information Model (Splunk)
CIMC	Cisco Integrated Management Controller
CNI	Container Network Interface
CNP / CCNP	CiliumNetworkPolicy / CiliumClusterwideNetworkPolicy
CRD	Custom Resource Definition (Kubernetes)
DCGM	NVIDIA Data Center GPU Manager
DCI	Data Center Interconnect
DLP	Data Loss Prevention
DMZ	Demilitarized Zone (network perimeter segment)

Term	Expansion
DPU	Data Processing Unit (NVIDIA BlueField)
eBGP	external Border Gateway Protocol
ESCU	Enterprise Security Content Update (Splunk)
EVPN	Ethernet VPN
FHFL	Full-Height, Full-Length (PCIe card form factor)
FIM	File Integrity Monitoring
GPU	Graphics Processing Unit
HA	High Availability
HEC	HTTP Event Collector (Splunk)
IPAM	IP Address Management
KV-cache	Key/Value attention cache (LLM inference)
LLM	Large Language Model
LVMS	OpenShift LVM Storage operator
MTU	Maximum Transmission Unit
NDFC	Cisco Nexus Dashboard Fabric Controller
NGC	NVIDIA GPU Cloud (registry)
NGFW	Next-Generation Firewall
NIM	NVIDIA Inference Microservice
NeMo	NVIDIA NeMo (model/retriever framework)
NVMe	Non-Volatile Memory Express (SSD interface)
OCP	OpenShift Container Platform
OOB	Out-of-Band (management)
OTel	OpenTelemetry
OVN	Open Virtual Network (default OpenShift CNI)
OWASP	Open Worldwide Application Security Project
PCIe	Peripheral Component Interconnect Express
PID	Product ID (orderable Cisco SKU)

Term	Expansion
PII	Personally Identifiable Information
RAG	Retrieval-Augmented Generation
RBA	Risk-Based Alerting (Splunk ES)
RBAC	Role-Based Access Control
RDMA	Remote Direct Memory Access
RHCOS	Red Hat CoreOS
RoCEv2	RDMA over Converged Ethernet v2
SaaS	Software as a Service
SCC	Cisco Security Cloud Control
SecOps	Security Operations
SKU	Stock Keeping Unit
SRE	Site Reliability Engineering
TC	Test Case (AI Defense test plan)
TLS	Transport Layer Security
ToR	Top-of-Rack (switch)
TopoLVM	Topology-aware LVM CSI driver
UCS	Cisco Unified Computing System
VIP	Virtual IP address
VLAN	Virtual LAN
VNI	VXLAN Network Identifier (L3VNI = Layer-3 VNI)
VRF	Virtual Routing and Forwarding
VXLAN	Virtual Extensible LAN
vLLM	High-throughput LLM inference/serving engine

Notation families used throughout – **Bx** ([trust boundaries](#)), **Dx** ([design decisions](#)), and **TC** (test cases) – are defined in the *Notation legend* at the front of this document. Known gaps and future work are consolidated in the [Known limitations and roadmap](#) section.

Appendix B: Validation summary matrix

This appendix groups the validation matrix by area so reviewers can scan each control, validation check, status, and source section without repeated area labels. Status legend:

- **Validated** = deployed and test-verified with evidence
- **Deployed** = deployed, verification pending
- **Partial** = deployed with a documented gap
- **Known issue** = deployed with a documented validation gap or unresolved result
- **Roadmap** = designed, not yet built

Network / fabric

Control / capability	Validation check	Status	Reference
VXLAN-EVPN leaf-spine underlay/overlay, NDFC-managed	Switches onboarded; ccStatus=In-Sync	Validated	Network and fabric design validation state
Enterprise-edge firewall (B1)	NGFW configured and enforcing	Roadmap – racked, unconfigured	Network and fabric design considerations
AI-factory perimeter (B2)	North-south fabric policy at border-leaf	Roadmap – temporary VLAN-3 bypass	Network and fabric design considerations
Cilium BGP LoadBalancer advertisement (host↔leaf eBGP)	Service VIP from 110.0.0.0/24 pool advertised over eBGP-multihop to both leaves (ASN 65500/65501 ↔ 65525, VRF saifvrf) and reachable	Validated	Application and workload security LoadBalancer IPAM and BGP

Compute / platform

Control / capability	Validation check	Status	Reference
UCS hardware (models, CPU, GPU, memory, firmware)	Live Cisco Intersight inventory	Validated	Compute, GPU, storage, and OpenShift components
OpenShift 4.19 clusters via Assisted Installer + Intersight	Both clusters installed, Cilium CNI at bootstrap	Validated	Compute, GPU, storage, and OpenShift design and implementation
GPU stack (GPU Operator, NIM/NIMCache/NIMService)	Models deployed and serving inference	Validated	Compute, GPU, storage, and OpenShift model deployment
Storage (LVMS/TopoLVM lvms-nvme, CephFS csi-fs-sc*)	Storage classes provisioned and bound	Validated	Compute, GPU, storage, and OpenShift storage classes

Control / capability	Validation check	Status	Reference
East-west 4×400 G SuperNIC fabric	Fabric configured / RDMA benchmarked	Roadmap – present, unvalidated	Compute, GPU, storage, and OpenShift east-west SuperNIC fabric

Workload security

Control / capability	Validation check	Status	Reference
Cilium CNI + identity-based policy (B5)	CNI running both clusters; identity policy enforced	Validated	App and workload security validation state
Default-deny per AI namespace (B4, D-5)	Positive/negative test: pod denied until policy added; lateral movement denied & logged	Pending validation – scenario defined; evidence package pending.	App and workload security design considerations
Tetragon base process/flow telemetry (B5)	Tetragon/Hubble events reach <code>cisco isovalent end-to-end</code>	Partial – enabled on both clusters; events reach Splunk from atl-ocp2. atl-ocp1 runtime-log export is intentionally out of scope for this phase to manage Splunk ingest capacity.	App and workload security validation state Security analytics design considerations
Tetragon TracingPolicy catalog (B5)	Catalog running as a standing deployment on each cluster	Partial – enabled on both clusters; path-all and base exec/flow active on atl-ocp2; full standing policy set pending on-cluster confirmation.	Kernel-level detection, App and workload security design considerations
Tetragon kernel enforcement + cryptominer AlertRule	Controlled lab trigger exercised the Sigkill/AlertRule path	Demonstrated in lab; publication evidence package pending.	App and workload security design considerations
Hubble flow visibility (Relay, UI, Timescape)	All four components deployed and exporting	Validated	Compute, GPU, storage, and OpenShift components

LLM security

Control / capability	Validation check	Status	Reference
AI Defense Hybrid Connector (on-premises data plane)	Data plane on-prem; only policy/metadata egress to SCC	Validated	LLM and model security validation state
Single-turn validation red-team (TC-01)	End-to-end run produces framework-mapped risk report	Validated (Pass)	LLM and model security validation state
Multi-turn adaptive red-team (TC-02)	Conversational jailbreak coverage	Partial – not supported on Hybrid Connector	LLM and model security design considerations

Control / capability	Validation check	Status	Reference
Inspection API – prompt injection (TC-03)	Benign allow vs. injection block	Validated (Pass)	LLM and model security validation state
Inspection API – DLP (TC-04)	Credential-shaped payload flagged	Known issue - credential-detector coverage gap documented; PII and response inspection confirmed working; remediation pending.	LLM and model security design considerations
Gateway – prompt injection / PII (TC-05)	Injection and SSN-shaped flows handled	Validated (Pass)	LLM and model security design considerations
Gateway – response guardrail (TC-06)	Unsafe model output blocked on return path	Validated (Pass)	LLM and model security design considerations
Indirect (RAG-borne) prompt injection	Post-retrieval augmented-prompt inspection	Roadmap	LLM and model security design considerations

Telemetry

Control / capability	Validation check	Status	Reference
Splunk OTel Collector, two-plane routing	Deployed both clusters; both exporters live	Partial – metrics plane live on both clusters; logs plane live on atl-ocp2; atl-ocp1 logs plane intentionally not enabled to bound Splunk ingest (capacity-driven scope)	Telemetry architecture operational and validation state
Metric allowlist + per-source pipelines	Named metrics reach Observability Cloud	Validated	Telemetry architecture metric allowlist
Token rotation runbook	Documented rotation procedure exists	Roadmap	Telemetry architecture design considerations

Analytics

Control / capability	Validation check	Status	Reference
Splunk ES + ESCU detections + RBA	Approximately 10 detections and five RBA rules enabled; CIM installed	Deployed	Security analytics operational and validation state
ES Findings / RBA firing end-to-end	Findings fire from live signals	Validated – Findings generated from live SAIF signals after RBA tuning.	Security analytics design considerations

Appendix C: Threat-model

This table maps enabled Splunk detections to the threat model. This appendix takes the design's point of view. For each adversary technique in scope (MITRE ATLAS and ATT&CK), it shows the enforcement control, the observation control, the owning trust boundary or document-local SAIF layer, and the current validation state. Layer labels such as SAIF L5 refer to this document's architecture layers, not OSI layers; SAIF L5 means AI Protection and SAIF L6 means AI Runtime. Together, those rows express the defense-in-depth principle of the solution architecture. Every technique is met by at least one enforcement control and one independent observation control.

Note: How to read the state column. "Validated" means deployed and test-verified with evidence. "Partial" and "pending" items are deployed with a specific, documented gap consolidated in Known limitations and roadmap. No row depends on a single control. For example, even where the default-deny negative test is still pending (T1567), the technique is constrained by AI Defense inspection points in the SAIF AI Protection layer and observed by Hubble in the SAIF Runtime Security layer. The defense-in-depth property holds by design even before every negative test is captured.

Table 25. SAIF security controls mapped to adversary techniques

Adversary technique	Enforcement control (design)	Observation control	Boundary / layer	State
LLM prompt injection (AML.T0051)	Cisco AI Defense prompt inspection on both surfaces – Inspection API and shared Gateway (D-6)	AI Defense verdict events > Splunk ES; ESCU high-volume-injection detection	SAIF L5 (AI Protection)	Validated (TC-03, TC-05)
Sensitive-information disclosure in model response (AML.T0048)	AI Defense response-side guardrail; DLP entity detectors on prompt and response	Response-side block events > Splunk ES	SAIF L5 (AI Protection)	Response guardrail validated (TC-06); DLP partial – named-PII entities caught, credential detector gap (TC-04, LLM and model security design considerations)
Model query abuse / cost harvesting (AML.T0034)	Single enforcement point at the shared LLM gateway (D-6); per-model GPU placement bounds blast radius	Gateway token/rate metrics (HEC-pushed) > Observability Cloud + ESCU token-burn detection	SAIF L5-L6 (AI Protection / AI Runtime)	Monitoring in place; automated throttle via Adaptive Response on the roadmap (Operations and lifecycle)
Data exfiltration to an external destination (T1567)	Cilium default-deny egress with identity/FQDN-based L7 policy (B4, D-1/D-5); nested boundaries B1-B5	Hubble flow logs + hubble drop /policy-verdict metrics; ESCU anomalous-outbound detection	SAIF L4 (Runtime Security) / B4	Policy deployed; default-deny negative and lateral-movement tests pending. See App and workload security design considerations
Malicious execution inside a workload pod (T1059)	Tetragon TracingPolicy in the kernel – exec, privileged-operation, and sensitive-file policies (B5); Cilium pod identity	Tetragon process/kprobe events > cisco isovalent > Splunk ES; ESCU privileged-file-access detection	SAIF L4 (Runtime Security) / B5	Runtime telemetry active and continuous on atl-ocp2; standing policy set and kernel-enforcement evidence pending

Adversary technique	Enforcement control (design)	Observation control	Boundary / layer	State
				(App and workload security design considerations , Known limitations and roadmap)

Table 26. Notation families used in this document

Family	Meaning	Example	Where defined
B1-B5 B=Boundary	Trust boundaries. The five nested boundaries across which the design changes its trust assumption and enforces a different control. Each boundary has both an enforcement control and an observation control.	B4 = the namespace boundary, enforced by CiliumNetworkPolicy, observed by Hubble flow logs	Five trust boundaries
SAIF L1-L7; L=Layer	Document-local SAIF architecture layers, not OSI layers. The stack runs from Network (SAIF L1) to AI Application (SAIF L7), with enforcement and observation responsibilities at each layer.	SAIF L4 = runtime security (Cilium/Hubble/Tetragon)	Seven layer model
D-1...D-n D=Design	Design decisions. The load-bearing decisions made for this build, each with its rationale, the alternatives considered, and any resulting constraint.	D-3 = data plane is 200 GbE Ethernet; RoCEv2 is not used	Bill of materials Validated design summary