

Sample Report

AI Infrastructure Security Assessment

A unified view across three assessments: one full adversarial threat report and two AI Security Posture Management snapshots, covering shadow AI, exposed inference, agentic-AI risk, secrets, supply chain, and regulatory exposure.

OVERALL RISK LEVEL — CRITICAL

74

HOSTS AFFECTED

46/100

AI-SPM POSTURE (LATEST)

51

CRITICAL + HIGH VIOLATIONS

3

EXTERNAL EXPOSURES

Prepared for: Sample — illustrative report

Assessments consolidated: Sample consolidated assessment · one threat report + two AI-SPM runs

Report generated: by ctrlAlx

Classification: ILLUSTRATIVE SAMPLE — synthetic data, no real client

01

Executive Dashboard

The organisation operates a large, largely unauthenticated AI/ML estate spanning OT/SCADA, Kubernetes, developer laptops and production servers. Confirmed internet-exposed inference APIs, hardcoded API keys in an internal platform repository, a persistence backdoor in an AI build image, and an agentic pipeline with an SQL-injection sink together place the environment at **critical** risk.

OVERALL RISK CRITICAL Driven by external exposure + secrets + backdoor	AI-SPM POSTURE (LATEST) 46/100 Rated POOR · 2026-04-20	POSTURE TREND ▼ 6 pts Worsened from 52 (2026-03-12)
UNIQUE HOSTS AFFECTED 74 Serving AI on sensitive ports	CONFIRMED EXTERNAL EXPOSURES 3 inference host (x2) + DMZ host	FRAMEWORKS FAILING EU AI Act 4 of 7 articles failing (latest run)

Findings by source assessment

Full Threat Report 2026-04-18 · XDR + repo + network hunts 902 findings · 38 critical · 84 high · 152 triaged / 74 hosts	AI-SPM Posture — Latest 2026-04-20 1,144 findings · 30 violations · 4 crit / 24 high / 2 med · EU AI Act FAIL	AI-SPM Posture — Baseline 2026-03-12 812 findings · 22 violations · 22 high · drift STABLE
---	--	---

Severity distribution (triaged threat items, n=152)



02

Critical Findings at a Glance

Seven cross-cutting issues account for the bulk of enterprise risk. Full evidence tables are in the Appendix.

Secrets in an internal platform repository	CRITICAL
Two live model-provider API keys are hardcoded in a test file inside the internal-platform repository. Present in git history; readable by anyone with repository access.	
Public-facing model API	CRITICAL
An inference host runs an unauthenticated vLLM/Python inference API on port 8000 with multiple bound addresses (incl. public-routable space) and confirmed external clients. Textbook path for model extraction and lateral pivot (ATLAS AML.T0024).	
Supply-chain backdoor	CRITICAL
A YARA persistence rule matched an AI build Dockerfile that installs cron/startup persistence and fetches a remote payload. Every CI/CD pipeline building this image would execute it.	
Agentic SQL-injection to RCE	CRITICAL
An agentic pipeline exposes an injection sink in a database-write tool allowing arbitrary SQL / command execution, plus ungoverned exfiltration via a search client and an embedding generator with no egress control.	
Unauthenticated inference sprawl	HIGH
74 hosts serve AI on sensitive ports (vLLM 8000, Ollama 11434, LM Studio 1234, llama.cpp 5001); 9 processes bind 0.0.0.0 as root. A vector database port-forwarded to all interfaces enables RAG data exfiltration.	
AI on OT / SCADA / POS	HIGH
An OT host serves a websocket on 8080 owned by a System process; point-of-sale terminals run llama.cpp on 5001; surveillance and mail systems own AI ports — port-squatting or covert-channel indicators.	
Shadow AI SaaS egress	HIGH
DNS telemetry shows dozens of hits to AI SaaS endpoints in 24h from non-approved hosts. An operational chatbot sends sensitive operational data to an external model provider.	

Posture trend (AI-SPM)

2026-03-12

52/100

→

2026-04-20

46/100

Posture regressed 6 points. The baseline showed partial failures across DORA, OWASP LLM, ATLAS and NIST; the latest run additionally fails **EU AI Act** (Articles 9, 12, 15, 72) and surfaces new agentic-AI attack surface.

03 Top Urgent Actions (Take Today)

Consolidated from the threat report and both AI-SPM runs, ordered by risk and blast radius. The full 16-item prioritised plan is in Appendix I.

- 1 Revoke the hardcoded model-provider API keys — today**
Revoke both keys via the provider console, reissue into a secrets manager / CI variables, then purge git history. Treat as fully compromised.
Owner: Security Engineering · Effort: ~1 hour
- 2 Isolate the exposed inference host from the internet**
Block port 8000 in/out at the perimeter. Bind the API to localhost or place behind an authenticated reverse proxy (mTLS / OAuth2). Enumerate loaded models and owner.
Owner: Network / Infrastructure · Effort: ~2 hours
- 3 Forensically investigate the Dockerfile backdoor**
Pull the Dockerfile, review the cron/startup code and payload URL, identify CI pipelines that built or ran the image, and disable those pipelines until cleared.
Owner: SOC / DevSecOps · Effort: 4–8 hours
- 4 Remediate the agentic SQL-injection sink + choke point**
Parameterise the database-write tool, split the prioritiser agent's read/write scope, and add human-in-the-loop on the subprocess runner.
Owner: Platform Development · Effort: ~9 hours
- 5 Fix 0.0.0.0 binds + add auth to inference endpoints**
Rebind 9 root-owned servers to localhost or a specific interface, strip --reload / debug logging, and stand up an API-gateway/auth layer for all inference instances.
Owner: DevOps / Platform · Effort: ~4 hours
- 6 Establish egress controls for AI data flows**
Gate the search-client and embedding egress, review the operational chatbot's data flow to the external provider, and extend AI-SaaS DNS monitoring to a 168-hour lookback.
Owner: Data Science / Legal / Security · Effort: ~1 day

04

Compliance & Regulatory Exposure

As a critical-infrastructure entity, the organisation is in scope for DORA and the EU AI Act. Note the divergence: automated posture controls return mostly PASS in the latest run, yet the adversarial threat report evidences substantial framework exposure. PASS reflects tooling coverage, not absence of risk.

Framework	Baseline (Mar 12)	Latest (Apr 20)	Notes
EU AI Act	0/7 PASS	4/7 FAIL	Latest run fails Art 9 (risk mgmt), 12 (logging), 15 (accuracy/robustness), 72 (post-market monitoring).
DORA	1/8 PARTIAL	0/8 PASS	Baseline flagged one control. Automated controls pass in latest run, but the threat report maps material ICT-risk exposure.
OWASP LLM Top 10	1/15 PARTIAL	0/15 PASS	Baseline flagged one item. Threat report independently evidences LLM01/02/03/06/07/09/10.
NIST AI RMF	1/11 PARTIAL	0/11 PASS	Baseline flagged GOVERN-1 (no AI asset registry / approval process).
MITRE ATLAS	1 PARTIAL	0 PASS	Baseline flagged one technique. Threat report maps 8 ATLAS techniques (see Appendix E).

Regulatory bottom line

HIGH

The latest AI-SPM run moves EU AI Act from PASS to FAIL (4/7). Combined with unauthenticated inference, ungoverned data egress to external model providers, and missing logging/monitoring, this is a credible basis for regulatory findings. Directory / identity integration is unconfigured, so identity attribution and access governance cannot currently be evidenced.

05

Threat Intelligence — Summary

Top mappings shown here; the complete MITRE ATLAS, ATT&CK, OWASP LLM and CVE matrices are in Appendices E–G.

MITRE ATLAS — highest-impact techniques

ATLAS ID	Technique	Evidence
AML.T0024	Exfiltration via ML Inference API	Inference host 8000 externally exposed; operational chatbot to external provider
AML.T0035	ML Supply Chain Compromise	Persistence YARA in build Dockerfile; agentic-framework wheel via internal proxy
AML.T0040	ML Model Backdoor	Hardcoded API keys + shadow AI endpoint in a test file
AML.T0006	Active Scanning	Multiple internal sources probing 3+ AI ports; top scanner hit 18 targets
AML.T0014	Discover ML Model Ontology	Unauthenticated Ollama 11434 model listing on multiple hosts

Exploitable CVEs in the estate (critical)

CVE	CVSS	Component	Applicability
CVE-2024-37032	9.8	Ollama path traversal / RCE	Ollama instances on 11434 — unpatched runtime
CVE-2024-0243	9.8	MLflow auth bypass	Model registry on 5000 — default configuration
CVE-2023-6975	9.8	MLflow path traversal → RCE	Same instance — artifact handler traversal

Plausible attack chain

CRITICAL

Recon sweep (AML.T0006) → reach the unauthenticated inference API from outside (T1190 / AML.T0024) → pivot to the vector database exposed via an all-interfaces port-forward → exfiltrate the RAG corpus and business logic. In parallel, a poisoned Dockerfile (AML.T0035) or a leaked provider key (T1552.001) gives persistence and cloud-side data egress.

A

AI Infrastructure — Critical Inference Hosts

All rows rated CRITICAL — inference services on sensitive ports. Host identifiers are synthetic. Source: endpoint telemetry.

Host (synthetic)	IP address	Component	Port	Process
inf-core-01	203.0.113.12 / 10.20.4.11	vLLM / Python inference API	8000	python3
ai-lab-01	10.30.34.10	Ollama LLM runtime	11434	ollama
kube-p04.corp.example	10.24.17.8	vLLM inference API	8000	python
kube-p05.corp.example	10.24.17.9	vLLM / inference manager	8000	manager
kube-p06.corp.example	10.24.17.10	vLLM / inference manager	8000	manager
svc-app-231	10.16.22.23	vLLM / generic inference API	8000	icman
svc-app-233	10.16.22.31	vLLM / generic inference API	8000	icman
svc-grc-101	10.16.22.71	vLLM / generic inference API	8000	icman
reg-node-a1	192.168.17.10	vLLM / generic inference API	8000	icman
reg-node-a2	192.168.17.13	vLLM / generic inference API	8000	icman
reg-node-b1	192.168.27.26	vLLM / generic inference API	8000	icman
reg-node-b2	192.168.27.28	vLLM / generic inference API	8000	icman
pos-term-7006	10.32.19.39	llama.cpp server	5001	pos-touch
pos-term-3006	10.32.21.35	llama.cpp server	5001	pos-touch
prism-server	192.168.36.10	llama.cpp server	5001	nginx
mbx-p04	10.24.3.4	Ollama	11434	mail-assistant
mbx-p03	10.24.3.3	Ollama + vLLM	11434 / 8000	edge-transport
ev-07	10.16.93.87	Ollama + alternate	11434 / 11435	dir-service
prod-vm	198.51.100.5 / 10.24.17.35	vLLM inference API	8000	python3
ws-p02.corp.example	10.24.205.26	vLLM / gunicorn inference	8000	gunicorn
code-host-516	10.11.232.135	vLLM inference API	8000	code-host
ml-mixed-005.corp.example	192.168.99.65	vLLM + text-gen-webui	8000 / 8080	node
scada-hmi	192.168.6.116	AI inference (text-gen-webui)	8080	system

Table truncated for the sample — the full estate lists 74 host groups.

B

AI Infrastructure — High-Severity Hosts

Web UI, notebook, registry and object-store exposures. Note the DMZ host and the very high connection counts on the notebook-hub proxy.

Host (synthetic)	IP	Component	Port	Process	Activity
web-ui-01	203.0.113.44	Open WebUI	3000	node	high connection volume
web-ui-02	10.24.15.25	Open WebUI	3000	node	internal
skw-p03	10.24.34.3	Open WebUI (docker-proxy)	3000	docker-proxy	internal
proxy-hub	192.168.111.115	Notebook hub proxy	8888	haproxy	very high conns / many clients
kube-p03.corp.example	10.24.43.148	Notebook hub	8081	hub	high conns
ml-report.corp.example	172.20.30.17	text-gen-webui	8080	java	internal
queue-mgmt	172.20.30.88	Model registry (MLflow)	5000	node	high conns
obj-store-01..06	10.24.17.11-16	Object storage (MinIO)	9000	minio	2,000+ conns each
web-dmz-01	link-local (DMZ)	AI API port (DMZ)	8081	system	EXTERNAL exposure

Note

HIGH

The DMZ host (web-dmz-01) is externally reachable, and the notebook-hub proxy shows a very high connection-to-client ratio — both warrant immediate review.

c

Anomalous Process–Port Pairings & Root Binds

A legitimate inference process would not own these ports; these indicate port reuse, process injection or port-squatting.

Host	Port	Expected AI component	Actual process	Risk
mbx-p03 / p04	11434	Ollama	mail-assistant / edge-transport	Mail service on Ollama port
evj-p01 / olg-p02	1234	LM Studio	monitoring-agent	Monitoring agent on LM Studio port
dlp-host-56	1234	LM Studio	dlp-agent	DLP agent on LM Studio port
pos-term-70xx	5001	llama.cpp	pos-touch	POS terminals on llama.cpp port
ev-07	11434	Ollama	storage-mgmt / dir-service	Storage/Dir services on Ollama port
vision-01	8001	Triton gRPC	node	Vision server – very high conns
scada-hmi	8080	text-gen-webui	system	SCADA System process on AI inference port
cam-node-457	8080	text-gen-webui	surveillance-suite	Surveillance software on AI port

Root-level binds & identity context

Directory enrichment returned no results — identity derived from process telemetry and mount metadata. Principals below are generic roles, not real users.

Principal (role)	Observed command / context	Assessment	Sev
ai-lab-01 / service-root	uvicorn --host 0.0.0.0 --port 8000 --reload (root)	Root AI server bound to all interfaces	CRITICAL
prod-host / service-root	uvicorn 0.0.0.0:8000 --reload (root)	Root AI server — production host	CRITICAL
oct-host / service-root	gunicorn 0.0.0.0:8000 --reload --log-level debug	Debug + root + all-interface bind	CRITICAL
kube-p01 / service-root	kubect1 port-forward svc/vectordb --address=0.0.0.0	Vector DB exposed — RAG exfil	CRITICAL
ci-runner / service-root	pip install (agentic framework) in CI	CI runner deploying agentic framework	HIGH
inf-core / service-root	uvicorn 0.0.0.0:8088 + orchestrator 0.0.0.0:8080	AI orchestration bound to all interfaces	CRITICAL
ml-workspace-owner	notebook PVC mount on kube-p06	Shared ML workspace mount	HIGH

D

External Exposure & Shadow-AI Egress

The primary inference host holds multiple bound addresses including public-routable space; its port-8000 inference API has confirmed external connectivity — directly exploitable for model extraction (AML.T0024).

Host	IP	Port	Process	External clients	Confirmed external IPs
inf-core-01	203.0.113.12	8000 (vLLM)	python3	2 ext. clients	203.0.113.1 (ext gateway)
inf-core-01	203.0.113.34	8000	python3	1 ext. client	198.51.100.3 (cloud NAT)
web-dmz-01	link-local (DMZ)	8081	system	multiple ext.	multiple external IPs

DNS / SaaS AI egress

Dozens of hits to AI SaaS in 24h from non-approved hosts. Provider names are generalised in this sample.

Sev	Domain (generalised)	Vendor	Hits (24h)	Source
CRITICAL	ai-saas-provider-1 (chat/LLM)	external	41	2 internal source IPs
CRITICAL	ai-saas-provider-1 (API mirror)	external	22	1 internal source IP
CRITICAL	ai-saas-provider-2 (LLM API)	external	2	2 internal source IPs

Supply-chain egress

HIGH

Remote-management agent traffic dominates outbound rows from a cluster of remote station hosts. An unauthorised AI meeting-recorder tool was also observed reaching an external recording service from a workstation.

E

Code Repository Findings

Repository scan partial (timed out) — a wider re-run is recommended. Paths below are generalised.

Repo	File	Rule	Sev	Detail
internal-platform	tests/test_embeddings.py	Hardcoded key (provider)	CRITICAL	Two model-provider keys hardcoded in a test file; live in git history, readable by anyone with repo access.
data-science	tools/build/Dockerfile	Persistence backdoor (YARA)	CRITICAL	Installs cron/startup persistence AND fetches a remote payload. Supply-chain: every CI/CD build executes it.
internal-platform	tests/test_embeddings.py	Shadow AI endpoint URL	HIGH	Hardcoded live external AI SaaS endpoint URL called from test code.

Blast radius

CRITICAL

The affected repository is an internal security-platform codebase — the security tooling itself. A leaked key here enables data exfiltration to an external provider, unbounded API cost, and tampering with the knowledge-base authoring pipeline.

LLM API calls in production / component code

Repo	File	Provider	Calls	Notes
internal-platform	orchestrator/runtime.py	Provider B	x2	Production runtime calling an external model
internal-platform	kb/authoring/generator.py	Provider A	x3	KB authoring pipeline
internal-platform	kb/embedder/client.py	Provider A	x1	Embedder client abstraction
data-science	chatbot/base.py	Provider A	x2	Operational chatbot to external model provider

Agentic stack

HIGH

An agent-orchestrator imports an agentic framework across many files; CI installs the framework and its dependencies. A model router routes to internal Ollama and external cloud — if Ollama is unauthenticated and model selection is caller-controlled, prompt injection or SSRF can redirect inference to attacker endpoints. Correlator and prioritiser agents call models inside security decision logic (OWASP LLM01).

F

Threat Mapping — MITRE ATLAS & ATT&CK

MITRE ATLAS

ATLAS ID	Technique	Evidence in findings
AML.T0000	Obtain Capabilities: AI Artifacts	Repo references to open models; model file downloads
AML.T0006	Active Scanning	Multiple sources scanning 3+ AI ports
AML.T0014	Discover ML Model Ontology	Unauthenticated Ollama model listing on multiple hosts
AML.T0024	Exfiltration via ML Inference API	Inference host 8000 external; operational chatbot to external provider
AML.T0035	ML Supply Chain Compromise	Persistence YARA in Dockerfile; framework wheel via internal proxy
AML.T0040	ML Model Backdoor	Hardcoded API keys in a test file; shadow AI endpoint
AML.T0043	Craft Adversarial Data	Correlator / prioritiser LLM calls with unsanitised input
AML.TA0002	ML Model Access	vLLM 8000 across nodes; Ollama on multiple hosts

MITRE ATT&CK

ATT&CK; ID	Technique	Evidence
T1059	Command & Scripting Interpreter	CI runner executing install chains via shell
T1053	Scheduled Task/Job (Cron)	YARA persistence backdoor in Dockerfile
T1552.001	Credentials in Files	Hardcoded provider keys in a test file
T1046	Network Service Scanning	Multiple IPs scanning 3+ AI ports
T1021.001	Remote Services (AI APIs)	uvicorn/gunicorn bound 0.0.0.0 on 9 hosts — lateral movement
T1071	Application Layer Protocol	LLM API to external providers over HTTPS — potential C2 channel
T1190	Exploit Public-Facing Application	Inference host externally exposed, no credentials
T1133	External Remote Services	Inference host 8000 internet-accessible; DMZ host 8081

G

Threat Mapping — OWASP LLM Top 10 & CVEs

OWASP LLM Top 10 (2025)

OWASP	Risk	Evidence
LLM01	Prompt Injection	Correlator / prioritiser ingest external alert data into an LLM
LLM02	Insecure Output Handling	Model router routes to multiple backends; output handling unverified
LLM03	Training Data / Supply Chain	YARA backdoor in Dockerfile; framework chain via internal proxy
LLM06	Excessive Agency	Multi-agent delegation; vector DB exposed via all-interface port-forward
LLM07	System Prompt Leakage	Open WebUI on multiple hosts (3000) — prompts may be extractable
LLM09	Misinformation	Operational chatbot used for operational safety decisions
LLM10	Model Theft	Unauthenticated vLLM via /v1/models + /v1/completions

CVE references (illustrative — real public CVEs, not the actual estate)

CVE	CVSS	Component	Applicability
CVE-2024-37032	9.8	Ollama path traversal / RCE	Ollama instances on 11434
CVE-2024-0243	9.8	MLflow auth bypass	Model registry on 5000 — default config
CVE-2023-6975	9.8	MLflow path traversal → RCE	Same instance — artifact handler
CVE-2024-27444	Medium	Notebook info-leak	Notebook hub series (8888)

H

Active Recon / Scanning Detected

Blue-team signal: internal sources probing multiple AI ports in patterns consistent with a Phase-2 discovery sweep. Correlate the top source against asset inventory.

Source IP	AI ports touched	Targets	Sev	Assessment
(broadcast)	5000,8000,8001,8080,8081,9000,11434 (11)	5	CRITICAL	Full AI port sweep — automated tool / adversary recon
10.24.23.2	8000,8080,8081,8888 (4)	18	CRITICAL	Broad internal sweep — 18 distinct AI hosts; not in inventory
10.24.17.8	8000,8080,8081,8888 (4)	9	HIGH	Node scanning sibling nodes
10.24.17.10	8000,8080,8081,8888 (4)	5	HIGH	Node scanning — health-check or adversarial
172.25.0.1	8000,8080,8081 (3)	4	HIGH	Internal scan

I Full Prioritised Recommendations

16 actions across Critical / High / Medium horizons.

Sev	Recommendation	Owner	Effort
CRITICAL	Rotate / revoke hardcoded provider API keys	Security Engineering	1 hour
CRITICAL	Isolate the inference host from external network	Network / Infra	2 hours
CRITICAL	Forensically investigate Dockerfile backdoor	SOC / DevSecOps	4-8 hours
CRITICAL	Firewall inference host 8000 + audit Python inference	Data Science / Infra	Half day
CRITICAL	Address 0.0.0.0 bind misconfigs on 9 servers	DevOps / Platform	4 hours
HIGH	Investigate SCADA host on AI inference port	OT Security	1 day
HIGH	Investigate anomalous process-port pairings	SOC / Windows Infra	1 day
HIGH	Protect model registry (MLflow) instance	DS Platform / Security	4 hours
HIGH	Secure notebook hub + port-forwarding	ML Platform / Security	Half day
HIGH	Investigate & block top scanner IP	SOC / Network	2 hours
HIGH	Remove AI inference ports from POS / OT systems	OT / Retail Security	Half day
HIGH	Secure orchestrator LLM key management	Platform Dev	2 days
HIGH	Assess operational chatbot external data flow	DS / Legal / Security	1 day
MEDIUM	API gateway / auth for all internal inference	Platform / DevOps	3-5 days
MEDIUM	Establish AI asset registry + approved-service policy	Security Architecture / CISO	1 week
MEDIUM	Reconfigure repo scan scope + secret scanning	DevSecOps	1 day

J

AI-SPM Violations — Latest Run (2026-04-20)

Posture 46/100 POOR · 1,144 findings · 30 violations (4 critical, 24 high, 2 medium). EU AI Act FAIL. An unattributed agent is flagged as an active threat actor.

Sev	Rule	Detail	Effort (h)	Risk
CRITICAL	Agent Injection Sink RCE/SQL/Shell	SQL injection in a database-write tool → arbitrary SQL execution	8	9
CRITICAL	Agent Choke Point Vulnerability	Prioritiser agent has excessive tool scope → privilege escalation	1	3
CRITICAL	Agent Tool Exfiltration Channel	Sensitive data leaves trust boundary via search client + embedding generator, no egress control	12	12
CRITICAL	Agent Sensitive Data Exfiltration	Same channel — sensitive data egress without control	14	9
HIGH	Agent Excessive Tool Scope	Prioritiser agent granted read+write access to incident data	4	2.7
HIGH	Agent Tool Boundary No HITL	Filter/parameter tampering in a subprocess runner; missing tool boundary	6	5.4
MEDIUM	Agent Trust Level Collapse	Prompt injection in a generic tool manipulates tool behaviour	3	1.8
MEDIUM	Agent Non-Repudiation Gap	Orchestrator handoff lacks reasoning-chain boundary	5	1.2

Remediation blueprint

CRITICAL

9 actions · ~54h · 2 quick wins. Order: choke point (1h) → injection sink (8h) → tool exfiltration (12h) → sensitive-data exfiltration (14h) → inference exposure (1h) → excessive tool scope (4h) → tool boundary/HITL (6h) → trust-level collapse (3h) → non-repudiation gap (5h).

K

Detection Signatures & Continuous Queries

What this activity looks like in XDR / SIEM.

Signal	What it indicates	Sev	Urgency
Connection to 8000/11434/5001/1234 from unexpected process	Port reuse / covert channel / malware	CRITICAL	Alert immediately
uvicorn/gunicorn/ollama launched with --host 0.0.0.0	AI service bound to all interfaces	HIGH	Alert + auto-ticket
kubectl port-forward --address=0.0.0.0	Internal service exposed to all interfaces	HIGH	Alert
Notebook API GET to 8888 without valid session cookie	Notebook enumeration	HIGH	Alert
Single IP hitting 3+ AI ports	Discovery sweep	CRITICAL	Alert
git push containing a provider key pattern	Provider key committed to git	CRITICAL	Alert + block push
New ollama/llamacpp server process on a workstation	Unsanctioned local LLM	HIGH	Alert
DNS to external AI providers from non-approved hosts	Shadow AI SaaS usage	MEDIUM	Log + review
cron / systemd modifications in Dockerfile during CI build	Supply-chain persistence attempt	CRITICAL	Block CI build

Recommended continuous queries (generic, tool-agnostic pseudocode)

```
-- AI processes binding to all interfaces
filter event=PROCESS and cmdline contains "0.0.0.0"
  and process in (uvicorn,gunicorn,ollama,vllm,python)

-- External client reaching AI inference ports
filter event=NETWORK and is_server=true
  and local_port in (8000,8001,11434,5001,1234)
  and remote_ip not in (RFC1918 ranges)
```

L

Scan Coverage & Methodology

Two coverage gaps materially limit identity attribution and shadow-AI visibility and should be closed before the next run.

Data source	Status	Notes
Endpoint telemetry	COMPLETE	Full findings across 74 hosts, 24h window
Network hunts	COMPLETE	All hunts returned; two throttled
Repository scan	PARTIAL	Timed out — re-run recommended
Directory / identity	NOT CONFIGURED	Enrichment returned no results — needs setup
DNS / CASB (AI SaaS)	PARTIAL	Possible proxy gap — extend to 168h lookback

WHAT WE COULD NOT SEE

This assessment maps roughly 80% of the AI estate. It does not, and cannot, see: personal devices on personal networks; on-premises systems with no egress that were out of telemetry scope; identity attribution where directory integration is unconfigured; and any system a user runs on a device outside managed coverage. We report what we can evidence, we state our confidence, and we name what we missed. A vendor claiming 100% coverage is claiming something no assessment can deliver.

Consolidation notes

This report merges three assessments: one full threat report and two AI-SPM posture runs. Where the posture runs and the threat report disagree on compliance status, both are shown — automated PASS reflects control coverage, not absence of risk.

Illustrative sample. This document mirrors the structure, methodology and layout of a real ctrlAlx consolidated assessment. Every organisation name, host identifier, IP address, repository path, username, finding count and CVE selection has been replaced with synthetic or generalised values. It represents no real client and contains no confidential data.