

TIRESIAS-ZT • THREAT & VULN INTEL

The governed threat-intel researcher.

A single agent that synthesizes the day's CVE, exploited-in-wild, and ATT&CK signal — and can only write what policy allows. It is the first node of ACME's SOC decision chain, and every one of its outputs is signed, tiered, and auditable.

● Asphodel + Tiresias-ZT • ACME Corporation • site
acme-site

THE STAKES

Threat intelligence is a trust problem before it is a data problem.

CVE feeds, exploited-in-wild advisories, and ATT&CK mappings arrive daily and feed decisions downstream. An ungoverned research agent can hallucinate a CVSS, cite a source that does not exist, or quietly write to a store it should never touch. In a SOC, that error propagates.

SIGNAL	Daily CVE / exploited-in-wild / ATT&CK synthesis The researcher's standing brief — the raw material the rest of the SOC reasons over	node 01
RISK	Unbounded model choice, unverifiable provenance, silent writes The failure modes governance is built to remove	deny-default
ANSWER	Every call tiered, every write signed, every source attributed Enforced at the proxy, not requested of the agent	enforced

THE WORKFORCE

One agent inside a governed fleet of 367.

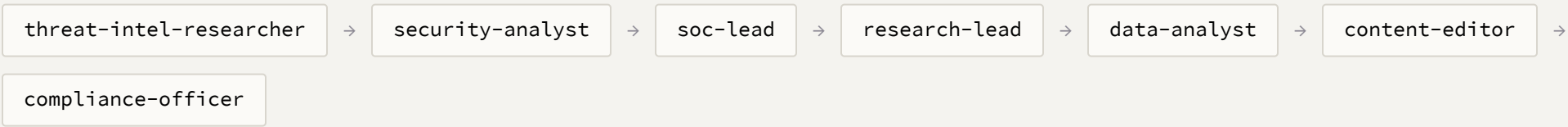
The threat-intel researcher is not a bespoke script. It is one tenant in a live zero-trust workforce running on GCP under a single deny-by-default policy — the same enforcement path every other agent traverses.

367 GOVERNED AGENTS	7 SHOWCASE / DAG AGENTS	360 FLEET AGENTS	325,154 DIRECTORY PRINCIPALS
---	---	--	--

Source: registry (367 agent tenants, postgres) · personas.json (7 showcase) · agent-manifest.json (360 fleet) · tzt_principal / Samba-AD — 120,120 humans + 205,034 NHI across 28 departments, 10 regions, 626 roles.

Node 01 of a 14-edge hand-off DAG.

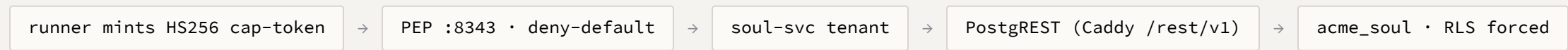
The 7 showcase agents form a topological hand-off DAG. The threat-intel researcher is where it begins — its synthesis is the input every downstream node depends on. Nothing reaches the SOC lead or compliance officer that did not start as a governed write here.



Source: `personas.json` — `container acme-runner-sched, runner.py schedule --interval 3600 (hourly). Fallback-to-floor on LLMError. Feeds the CISO and ATT&CK mindsets downstream.`

A research write never touches the data store directly.

When the researcher writes its daily brief, the call is minted, validated, forwarded verbatim only on allow, and signed into an audit chain. Every hop below is enforced infrastructure, not agent goodwill.



| Allow forwards verbatim. Deny is the default. The agent cannot opt out of the proxy.

Source: Governed request flow – PEP validates X-Tiresias-API-Key + X-Capability-Token, policy eval; signs the call into a SoulKey Ed25519 audit chain → CP. Usage → /cp/usage/report → tzt_usage → Grafana.

LIVE PROOF

Real numbers off the running instance.

This is not a mock. The fleet runs on independent rolling batches against a live Cloud SQL data tier, and the economics come straight off the ROI board.

1–2 GOVERNED JOBS / DAY / AGENT	6 FLEET_BATCH SIZE	10h MIN GAP BETWEEN JOBS	9 MODEL TIERS · 3 PROVIDERS
---	--	--	---

Source: `fleet_runner.py` – `FLEET_BATCH=6`, `INTERVAL=300s`, `MIN_GAP=10h`. `config.json` – 9 tiers across OpenRouter, Ollama Cloud, HuggingFace. `tzt_usage` → Grafana uid `acme-roi-economics`.

Homed on free compute, with a ceiling it can reach when it matters.

Every agent carries a home / ceiling / floor tier. The threat-intel researcher homes on `ollama_oss` — free self-hosted gpt-oss:120b — so routine daily synthesis costs effectively nothing. Its floor is the same free tier, giving a reliable degrade path.

HOME TIER

ollama_oss — gpt-oss:120b on Ollama Cloud. Routine CVE/ATT&CK synthesis routes here first. Routing favors free/Ollama by design.

FLOOR TIER

ollama_oss — free-homed floor is ollama_oss for a reliable degrade. On `LLMError` the runner falls back to floor rather than failing the job.

Source: `fleet_runner.py` routing — home free_large/ollama_oss/ollama_deep; free-homed floor = ollama_oss. Fallback-to-floor on `LLMError`.

Daily brief → tiered model → governed CoT + memory write.

One job is one full governed cycle. The researcher takes its domain brief, calls a tiered model, and commits reasoning and findings through the enforcement path — nothing lands in acme_soul that skipped the proxy.

01	Domain brief — daily CVE / exploited-in-wild / ATT&CK synthesis The standing task for node 01 of the SOC DAG	input
02	Tiered LLM call via per-agent capability token Home tier ollama_oss; PEP validates before any forward	POST /v1/cot/capture
03	Governed chain-of-thought + memory write ContentInspection enforces source_type + authored_by	POST /v1/memory/write
04	Usage reported, call signed into audit chain Downstream nodes read via /v1/memory/search	/cp/usage/report

Source: `fleet_runner.py` job shape — domain brief → tiered LLM → governed CoT + memory write → usage report. `soul-svc ContentInspection (enforce)` requires `source_type + authored_by`.

When a keyword changes the stakes, the tier moves with it.

A routine advisory and an active breach do not deserve the same model. Keyword escalation lets the researcher lift off its free home tier toward its ceiling when the signal warrants — automatically, within policy bounds.

TRIGGER KEYWORDS

`+incident +breach +critical`

Presence in the brief escalates the call above the home tier toward the ceiling for that job.

BOUNDED, NOT UNBOUNDED

Escalation is capped by the agent's ceiling tier. The researcher cannot silently jump to opus on a whim — the tier ladder is policy, and the PEP still governs the call.

Source: `personas.json` — per-agent home/ceiling/floor tiers + keyword escalation (`+incident / +breach / +critical` → `escalate`).

Signed, isolated, and watched for drift.

Governance is not a report generated after the fact — it is the write path itself.
Every research output is attributed, tenant-isolated at the database, and
monitored for behavioral drift.

SIGN	SoulKey Ed25519 audit chain, durable WAL at /data/audit-wal Every governed call signed and shipped to the control plane	Ed25519
ISOLATE	Per-agent soul-svc tenant, FORCE RLS on acme_soul Cross-tenant reads require SET ROLE service_role	RLS forced
ATTRIBUTE	ContentInspection requires source_type + authored_by No unattributed research can be written	enforce
WATCH	SoulWatch behavioral drift → auto-quarantine SIEM: RFC5424 syslog → Promtail:1514 → Loki	quarantine

Source: Enforcement Plane — signed audit chain, FORCE RLS, ContentInspection, SoulWatch drift → auto-quarantine (shared postgres store).

ECONOMICS

The self-host-vs-API answer, per 1M tokens.

Homing the researcher on free self-hosted compute is not a compromise — it is the economically dominant choice for high-volume daily synthesis. The ROI board makes the gap explicit.

MODEL / TIER	\$ / CALL (OBSERVED)	RELATIVE COST
opus-4.8 (paid ceiling)	\$0.19	
sonnet-4.5 (paid)	\$0.03	
deepseek-v4-pro	\$0.002	
gpt-oss:120b (ollama_oss · home)	\$0.0015	
Llama-70B (HuggingFace)	\$0.0012	
gemma-4-31b:free	\$0	

Source: Grafana ROI dashboard uid acme-roi-economics — observed \$/call. Free-tier saturation (hermes-405b:free chronic 429) is the self-host case.

CLOSE

Governed at the source, trusted downstream.

The threat-intel researcher proves the model: an agent doing real work — daily CVE, exploited-in-wild, and ATT&CK synthesis — where every output is tiered for cost, signed for provenance, isolated by tenant, and watched for drift. Node 01 is trustworthy, so the SOC chain that reads from it can be too.

A governed write here is what makes the CISO and ATT&CK mindsets downstream worth believing.

● Asphodel + Tiresias-ZT • ACME Corporation • console.asphodel.ai •
pdp.tiresias.watch