

CephSentinel AI

FIRST PILOT AVAILABLE

Continuous monitoring and operational automation for Ceph fleets

Module of the SentinelForge AI ecosystem · on-premise · sovereign · open-source · AI-native

The problem

Ceph is reliable but it is not simple. A cluster in production speaks in thresholds and status codes: it says that something is wrong, not why it is happening or what to do. Those who manage Ceph spend their time translating warnings into decisions — and that translation lives in the heads of two or three people.

With more than one cluster the problem changes in nature. Clusters belonging to different customers, in different data centers, some behind firewalls that expose nothing. No overall view: to know where to look today you have to open one cluster at a time. And the most costly errors are not the failures, they are the oversights: the noout flag left on for weeks, the maintenance window that nobody closed, the critical warning ignored for hours because nobody was watching the dashboard.

The solution

CephSentinel AI is the module of the SentinelForge AI ecosystem dedicated to Ceph fleets. It is not one more alerting system: it is a daemon that observes autonomously, correlates, deduplicates, classifies and prioritizes signals instead of waiting for someone to ask a question, and manages several clusters at the same time — including geographically separate ones that are not co-located with the system itself.

Every managed cluster is a first-class object, with dedicated credentials, its own contacts and a configurable automation policy: one customer may authorize more autonomy than another on the same operation. Metrics, events, forecasts and audit logs are partitioned per cluster right from the data schema.

The heart of it is neither the chat nor the model: it is the fleet view. A single screen that answers "where should I look today", with the clusters sorted by urgency. It is what an operator with dozens of clusters does not have today — and which no traditional tool offers.

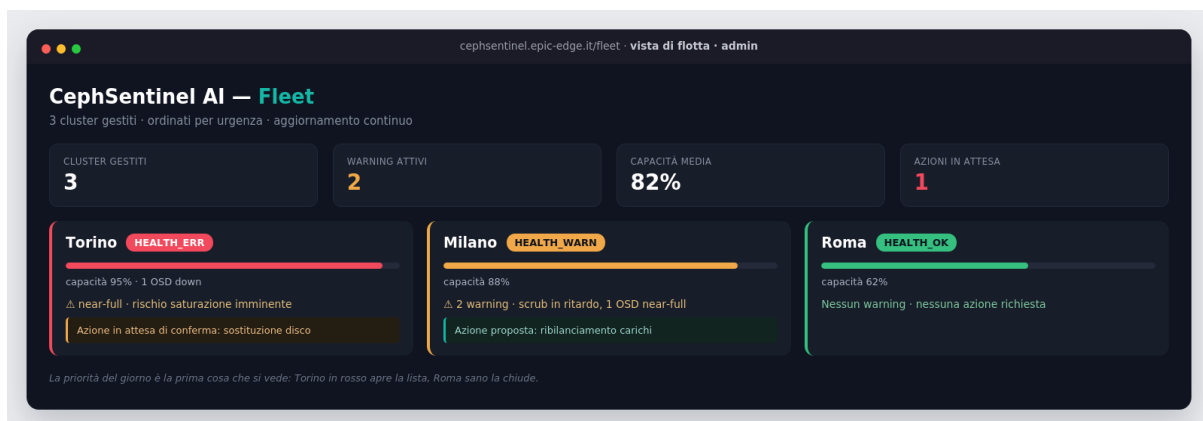


Figure 1 — Fleet view: the clusters sorted by urgency. The priority of the day is the first thing you see.

The model proposes. The automation engine executes.

How it works

1 · OBSERVE

Collector per cluster: structured status, topology, usage per pool, placement group status, SMART metrics. Continuous cycle plus immediate trigger on critical warnings.

2 · INTERPRET

Rule-based layer for known thresholds. LLM layer only where context is genuinely needed: a natural-language explanation tailored to that cluster.

3 · CORRELATE AND PREDICT

It connects signals that individual thresholds would never connect — the network problem on a rack and the OSD flapping on that rack are a single incident. Capacity projections based on the historical trend.

4 · PROPOSE AND ACT

A readable action plan, classified by risk according to the policy of that cluster. Traced execution, never the model directly on the cluster.

Anatomia di un incidente: dai segnali alla causa

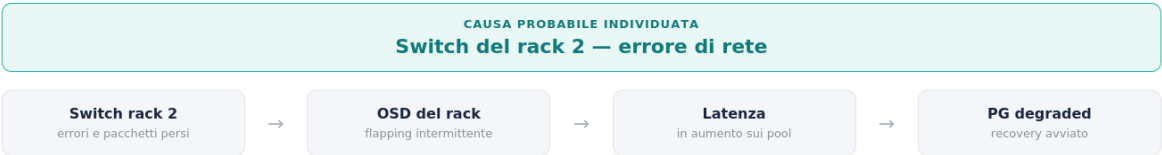
CephSentinel raccoglie segnali da fonti diverse e li correla in una diagnosi — prima che diventino disservizio.



Figure 2 — Anatomy of an incident: different signals over time correlated into a single diagnosis.

Un solo incidente, non quattro allarmi

L'analisi delle cause collega segnali che le soglie singole non collegherebbero mai.



Quattro allarmi separati su una dashboard tradizionale; qui, un unico incidente con una causa a monte. È dove l'AI aggiunge valore rispetto a un alerting per soglie.

Figure 3 — Root cause analysis: four separate alarms, a single upstream cause.

Autonomy proportionate to risk

● **AUTONOMOUS** Collection, diagnostics, notification

Extended health checks, daemon version verification, delayed scrubs, quorum, disk metric scraping, turning on the LED of the disk to be replaced, dry-runs and previews. No impact on data.

● **SUPERVISED** Day-to-day operations, with confirmation

Controlled OSD removal with impact estimation, end-to-end disk replacement, host maintenance mode, orchestration of cluster expansion (hosts and OSDs), phased rolling upgrades, recovery throttling between peak hours and the maintenance window, temporary flags with automatic timeout, pool and credential management.

● **ADVISORY ONLY** The system explains, the operator decides

Changes to the CRUSH map, fill thresholds, forced host removal, purge. The system shows the impact simulation and the guided steps; the final command always remains in the operator's hands.

Autonomia proporzionale al rischio

Ogni operazione ricade in uno di tre livelli, configurabili per cluster e per cliente.



Si parte con tutto supervisionato; l'autonomia si allarga con la fiducia, cluster per cluster.

Figure 4 — The three levels of autonomy, configurable per cluster and per customer.

Key features

- **Fleet view, not cluster view** — the home page answers “where should I look today”: clusters sorted by urgency, active warnings, actions awaiting confirmation, and a cluster that has lost connectivity for hours is flagged differently from a healthy one.
- **Real health, not just the Ceph code** — a cluster that is technically HEALTH_OK but with a forecast of imminent capacity exhaustion does not appear “all green”.
- **Two connectivity models** — direct polling towards the customer's endpoint, or a lightweight agent installed at the customer's site that connects outbound: no inbound port to open, a requirement that is often non-negotiable in storage-critical environments.
- **No constraint on the deployment method** — it operates on any Ceph cluster via CLI, or via API where available: not only clusters managed by cephadm.
- **Two-layer engine** — the deterministic rules cover known cases with minimal latency and zero risk of misinterpretation; the LLM steps in only where context is needed. A model is not invoked for a threshold check.
- **Multi-signal correlation** — the part where AI adds real value compared with traditional alerting: two distinct warnings that are in reality a single incident with one root cause.

- **The model never executes commands on the cluster** — it generates a structured plan; execution goes through a separate engine (AWX + Ansible). It is this decoupling that makes audit, dry-run and revocation before execution possible.
- **Full traceability** — every playbook executed is versioned on Git and explicitly linked to the target cluster and to the signal that generated it. What was done and why is always open to inspection.
- **Isolation per customer** — dedicated credentials per cluster, never shared: a compromised cluster gives no access to the others.
- **Automatic escalation** — if a high-severity supervised action remains unconfirmed beyond the threshold, the system raises the notification level. A critical warning ignored for hours is in itself a failure, and must be prevented structurally.
- **On-premise LLM** — no cluster data, no topology, no credential leaves the perimeter.

Benefits for the customer

One fleet, one operator

Managing ten clusters no longer means ten open sessions. The priority of the day is the first thing you see.

Fewer costly oversights

Flags with timeouts, supervised maintenance windows, escalation on missed confirmations: the most common operational errors become structurally difficult.

Distributed Ceph expertise

The contextualized explanation makes those without twenty years of Ceph behind them operational, without taking control away from those who have.

Every action defensible in an audit

What was executed, on which cluster, for which signal, with which approval. A requirement, not an extra, in regulated contexts.

Autonomy tailored to the customer

Each customer explicitly authorizes what they want to automate. You start cautiously and widen the scope as trust grows.

Data sovereignty preserved

On-premise engine: monitoring with AI does not require sending the telemetry of your own clusters to an external service.

Stack and technologies

Telemetry	Prometheus + native ceph-mgr exporter — the Ceph standard, not reinvented
Collector	Dedicated Python service per cluster — complete structured status, not just numerical series
LLM engine	Open-weight models run on-premise via Ollama
Execution	AWX + Ansible, on top of the native Ceph orchestrator where present; via CLI or API on clusters that do not expose it
Versioning and audit	GitLab — every playbook linked to its cluster and originating event

Middleware	FastAPI — REST for status and confirmations, WebSocket/SSE for real-time push
Dashboard	React — fleet view and per-cluster drill-down
Notifications	Email, Slack, Teams, SMS — configurable per cluster

Who it is for

- **Service providers and MSPs** — those who manage Ceph on behalf of third-party customers, with different contacts, SLAs and autonomy levels for each.
- **Telco & ISP** — several clusters in several data centers, where an overall view does not exist today.
- **Multi-data-center enterprises** — storage distributed across different sites, with a central team that cannot oversee everything by hand.
- **Public administration** — data sovereignty and traceability requirements for infrastructure interventions.
- **Finance & Banking** — high-availability storage where every change must leave an audit trail.

Deployment model

Status	First pilot available — multi-cluster observation and supervised actions
Mode	On-premise, within the customer's perimeter
Connectivity	Direct polling towards the cluster endpoint, or a lightweight outbound agent for clusters behind restrictive firewalls
Compatibility	Ceph clusters in production, regardless of the deployment method — access via CLI or API
Onboarding	Cluster registration through a guided wizard with an immediate reachability test; a cautious initial policy, refinable later
Autonomy	Configurable per cluster across the three risk levels; the scope of the automations activated is agreed at activation — everything starts supervised

The Epic Edge ecosystem



Every Epic Edge product is part of SentinelForge AI: distinct modules on the same substrate — on-premise LLM, traced execution, autonomy policies, audit, multi-provider middleware.

Architectural principles

Building an AI automation platform for infrastructure means holding together things that usually live separately: integrating technologically different providers, maintaining independent adapters that evolve without touching the rest, guaranteeing a complete audit of every action, preserving data sovereignty and — above all — separating the engine that decides from the engine that executes.

This is why CephSentinel AI was created as a module of the SentinelForge AI ecosystem: not a function added to an existing tool, but an architecture designed by people who have lived with those constraints in production.

Why Epic Edge

Ceph in production, not in theory

Real clusters, real incidents, real SLAs. The expertise comes from the field.

Patterns already proven

LLM engine, execution via AWX+Git, real-time streaming: already validated in production with EdgeForge AI.

Open-source and sovereign

No lock-in, no telemetry to third parties, everything inside the customer's perimeter.

How many Ceph clusters are you managing by hand?

The first pilot is available. Let us register one of your clusters and show you what it sees — without touching anything, until you decide otherwise.

andrea.martra@epic-edge.it · www.epic-edge.it/sentinelforgeai



Epic Edge S.r.l. · Via Michele Buniva, 26 — 10064 Pinerolo (TO), Italy · info@epic-edge.it · www.epic-edge.it

OPEN-SOURCE · SOVEREIGN · ENTERPRISE-GRADE · AI-DRIVEN