

SentinelForge AI

ECOSISTEMA AI-OPS

L'ecosistema AI per le operazioni infrastrutturali — dal provisioning all'ottimizzazione

FORGE · PROTECT · OPTIMIZE — on-premise · sovrano · open-source



Figura 1 — L'ecosistema SentinelForge AI: tre moduli (Forge, Protect, Optimize) sullo stesso substrato — motore LLM on-premise, esecuzione tracciata, policy di autonomia, audit, middleware multi-provider.

Il problema

Gestire un'infrastruttura cloud richiede tre competenze distinte in tre momenti diversi: costruire gli ambienti, tenerli in salute, ottimizzarne i consumi. Oggi ciascuna di queste fasi vive in strumenti separati, con conoscenza concentrata in poche persone e processi che non parlano tra loro.

A questo si aggiunge un vincolo che nel settore pubblico e regolamentato è spesso non negoziabile: le soluzioni di automazione AI disponibili sul mercato impongono di far uscire dati, telemetria e credenziali verso servizi esterni. Automazione intelligente e sovranità del dato, oggi, sembrano alternative.

La soluzione

SentinelForge AI è l'ecosistema AI di Epic Edge per l'intero ciclo di vita dell'infrastruttura. Non tre applicazioni scollegate, ma moduli che condividono lo stesso substrato — motore linguistico on-premise, esecuzione deterministica tracciata, policy di autonomia, audit — e coprono le tre fasi: costruire (Forge), sorvegliare (Protect), ottimizzare (Optimize).

Si adotta un modulo alla volta, sull'infrastruttura che già si ha; il valore cresce quando i moduli lavorano insieme sulla stessa fondazione. In ogni fase l'intelligenza artificiale opera interamente nel perimetro dell'organizzazione: nessun dato lascia la rete.

Il modello propone. Il motore di automazione esegue.

I tre moduli

EdgeForge AI FORGE DISPONIBILE

Provisioning di infrastruttura in linguaggio naturale, su OpenStack e Proxmox VE. L'utente descrive ciò che serve; l'assistente progetta l'architettura, dimensiona le risorse, ne stima il costo e — dopo l'approvazione — esegue il deploy in modo automatizzato e tracciato.

CephSentinel AI PROTECT PRIMO PILOT

Monitoring continuo e automazione operativa per flotte Ceph, anche multi-cluster e su cluster remoti. Correla, deduplica e priorizza i segnali, e propone — eseguendole su conferma — le operazioni quotidiane, con tre livelli di autonomia configurabili per cluster.

CloudSentinel AI OPTIMIZE IN SVILUPPO · 2027

Sorveglianza intelligente e ottimizzazione energetica sull'infrastruttura fisica: anomaly detection, analisi automatica delle cause e consolidamento dei carichi per ridurre i consumi. In sviluppo; raccolta dati di esercizio già avviata.

Il substrato comune

I moduli non sono prodotti separati con tecnologie diverse: poggiano tutti sulla stessa fondazione. È questo a rendere l'ecosistema più della somma delle sue parti — e a permettere di aggiungere nuovi domini senza ricostruire ogni volta da zero.

- **Motore linguistico on-premise** — modelli open-weight eseguiti nel perimetro del cliente (Ollama): nessun dato, nessuna telemetria, nessuna credenziale lascia la rete.
- **Il modello non esegue mai** — genera un piano; un motore di esecuzione deterministico separato (AWX + Ansible) lo esegue. È ciò che rende possibili audit, dry-run e revoca prima dell'esecuzione.
- **Policy di autonomia** — gli stessi tre livelli in tutto l'ecosistema: segnalazione, esecuzione previa approvazione, esecuzione automatica limitata alle azioni a basso rischio.
- **Tracciabilità nativa** — ogni playbook è versionato su Git e legato al segnale che l'ha generato: cosa è stato fatto, dove e perché, resta sempre ispezionabile.
- **Astrazione multi-provider** — un livello intermedio normalizza le tecnologie sottostanti: aggiungere un dominio (OpenStack, Ceph, Kubernetes, Proxmox...) significa estendere la piattaforma, non riscriverla.

Caratteristiche chiave

- **Sovrano by design** — l'automazione AI non impone di esporre dati e infrastruttura a servizi extraeuropei.

- **Modulare** — si adotta il modulo che serve, quando serve; partire da uno non preclude gli altri.
- **Coerente** — stessa logica di approvazione, autonomia e audit in ogni modulo: si impara una volta, vale per tutto l’ecosistema.
- **Open-source alla base** — costruito sugli stessi stack che alimentano i cloud sovrani e telco europei: nessun lock-in proprietario.
- **Difendibile in audit** — ogni azione, manuale o assistita, lascia una traccia completa: requisito, non accessorio, nei contesti regolamentati.

Benefici per il cliente

Un ecosistema, non tre strumenti

Costruire, sorvegliare e ottimizzare vivono nella stessa piattaforma, con la stessa esperienza — non in silos separati con competenze diverse.

Sovranità preservata

L’intelligenza artificiale entra nell’operatività senza rinunciare al controllo del dato: tutto resta nel perimetro.

Competenza distribuita

La conoscenza operativa è nel sistema, non in poche persone: onboarding più rapido e meno dipendenza da profili rari.

Adozione a basso rischio

Si parte da un modulo su un ambiente reale, con tutto supervisionato; l’autonomia si allarga con la fiducia.

Stack e tecnologie

Domini infrastrutturali	OpenStack · Ceph · Kubernetes · Proxmox VE
Motore LLM	Modelli open-weight eseguiti on-premise via Ollama
Esecuzione e automazione	AWX + Ansible, con i playbook versionati su GitLab
Osservabilità	Prometheus e Grafana — valorizzando gli strumenti già presenti
Identità e accesso	Keycloak — SSO e federazione con le directory aziendali

A chi si rivolge

- **Telco & ISP** — operatori con infrastruttura open-source di produzione e volumi elevati di operazioni ripetitive.
- **Cloud & Service provider** — chi eroga servizi cloud e cerca di ridurre il costo operativo lungo l’intero ciclo di vita.
- **Pubblica Amministrazione** — automazione AI compatibile con i requisiti di sovranità del dato e di tracciabilità.
- **Finance & Banking** — contesti dove ogni intervento sull’infrastruttura deve lasciare una traccia di audit.

Stato e adozione

EdgeForge AI (Forge)	Disponibile — provisioning operativo su OpenStack e Proxmox VE
CephSentinel AI (Protect)	Primo pilot disponibile — osservazione multi-cluster e azioni supervisionate
CloudSentinel AI (Optimize)	In sviluppo — roadmap 2027; raccolta dati di esercizio avviata
Modello di adozione	Un modulo alla volta, sullo stesso substrato: partire da uno non preclude gli altri

Perché è diverso da un'AI gestita in cloud

La differenza non è nel numero di funzioni: è in dove vivono i dati, il modello e il controllo.

	Epic Edge SentinelForge AI	AI-as-a-Service gestita in cloud
AI eseguita on-premise, nel perimetro del cliente	● Sì	● No
Nessun dato o telemetria inviati all'esterno	● Sì	● No
Stack open-source, senza lock-in proprietario	● Sì	● No tecnologia proprietaria
Audit completo e tracciabile di ogni azione	● Sì Git + audit trail	● Parziale log applicativi
Sovranità e giurisdizione del dato	● Sì	● Dipende dal provider e dalla region
White-label per l'operatore nazionale	● Sì in rilascio	● No
Avvio immediato senza infrastruttura propria	● No richiede il proprio ambiente	● Sì

La sovranità ha un prezzo, ed è dichiarato: serve un'infrastruttura propria. In cambio, dati, modello e controllo non lasciano mai il perimetro.

Figura 2 — Epic Edge vs. AI gestita in cloud: la differenza è in dove vivono dati, modello e controllo. La sovranità richiede un'infrastruttura propria; in cambio nulla lascia il perimetro.

Roadmap di prodotto

Dove sta andando la piattaforma — con lo stato reale di ciascun elemento, non promesse.

Disponibile in produzione e in pilot	Roadmap in sviluppo - 2027	Visione strategica direzione della piattaforma
● EdgeForge AI provisioning in linguaggio naturale — live ● CephSentinel AI operazioni sulle flotte Ceph — primo pilot ● CMP — portale unificato multi-provider — in implementazione ● OpenStack · Proxmox VE provisioning operativo ● Osservabilità Grafana metriche per ambiente — disponibile	● CloudSentinel AI sorveglianza predittiva ed efficienza energetica ● VMware provisioning — in integrazione ● Capacity planning cross-provider vista aggregata nel portale ● White-label completo per l'operatore nazionale ● Agente CephSentinel per cluster dietro firewall	● AWS · Azure gestiti governo degli estate hyperscaler ● Osservabilità estesa oltre le metriche: log e tracce ● Operazioni autonome il cloud che si gestisce da sé, sotto controllo
● disponibile ● pilot / in implementazione ● pianificato 2027 ● visione		

Figura 3 — La roadmap di prodotto: disponibile oggi, pianificato 2027, direzione della piattaforma — ogni voce con il suo stato reale.

Principi architetturali

Realizzare una piattaforma di automazione AI per l'infrastruttura significa tenere insieme cose che di solito vivono separate: integrare provider tecnologicamente diversi, mantenere adapter indipendenti che evolvono senza toccare il resto, garantire un audit completo di ogni

azione, preservare la sovranità del dato e — soprattutto — separare il motore che decide dal motore che esegue.

È per questo che SentinelForge AI è stato progettato come ecosistema modulare: non una collezione di funzioni, ma un'architettura pensata da chi quei vincoli li ha vissuti in produzione.

Perché Epic Edge

25+ anni di operations

Infrastrutture di produzione reali, non concetti su slide. La competenza viene dal campo.

Già in produzione

EdgeForge AI è live e CephSentinel AI in pilot: l'ecosistema non è una roadmap, metà già gira.

Sovrano e open-source

AI on-premise, stack aperto, nessuna dipendenza da servizi extraeuropei.

Da quale fase volete partire?

Forge, Protect o Optimize: identifichiamo il modulo con il ritorno più rapido sul vostro ambiente e partiamo da lì, con un pilot a rischio contenuto.

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

EdgeForge AI

LIVE

CephSentinel AI

PILOT

CMP

IN IMPLEMENTAZIONE

CloudSentinel AI

ROADMAP

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

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