

Cloud Management Platform

Piattaforma di Cloud Management AI-native — multi-provider, sovrana e white-label

Documento di presentazione del prodotto

Versione 1.9.4 · Luglio 2026

Scopo del documento

Questo documento descrive, in una prospettiva di prodotto e di servizio, la Cloud Management Platform sviluppata da Epic Edge: a chi si rivolge, quali esigenze risponde, come si presenta agli utenti e quale valore porta a un contesto di cloud nazionale e di pubblica amministrazione. Non contiene dettagli implementativi o descrizioni architetturali di livello tecnico, che sono oggetto di documentazione separata.

1. Sintesi

La Cloud Management Platform (CMP) di Epic Edge è il punto di accesso unico attraverso il quale un'organizzazione governa le proprie risorse cloud, indipendentemente dalla tecnologia su cui queste risorse sono realizzate.

Oggi la maggior parte delle organizzazioni non ha un solo cloud: ha una piattaforma privata basata su tecnologie open source, un ambiente di virtualizzazione tradizionale, talvolta risorse presso operatori pubblici internazionali. Ognuno di questi ambienti ha una propria console, un proprio linguaggio, una propria logica di accesso e di misurazione dei consumi. Il risultato è una frammentazione che genera costi di formazione, errori operativi, difficoltà di rendicontazione e, soprattutto, perdita di controllo complessivo.

La CMP di Epic Edge risolve questa frammentazione presentando tutti gli ambienti attraverso un'unica interfaccia coerente, con un unico accesso, un'unica logica di visualizzazione dei consumi e due componenti di intelligenza artificiale integrate: EdgeForge AI, che esegue le operazioni richieste dall'utente in linguaggio naturale, e CloudSentinel AI, che sorveglia lo stato dell'infrastruttura, ne anticipa i guasti e ne ottimizza i consumi energetici. La piattaforma è interamente basata su tecnologie open source, è installabile presso il cliente e non prevede alcuna dipendenza da servizi extraeuropei.

In una frase

Un'unica piattaforma, personalizzabile con il marchio dell'operatore che la adotta, che rende governabili in modo semplice e uniforme ambienti cloud diversi tra loro, affiancando all'utente un assistente in linguaggio naturale per le operazioni quotidiane e all'operatore un sistema intelligente di sorveglianza e ottimizzazione.

2. Il contesto e le esigenze a cui risponde

Il progetto nasce dall'esperienza diretta di Epic Edge nella progettazione e gestione di infrastrutture cloud per clienti enterprise e per operatori di servizi. Da questa esperienza emergono con costanza quattro esigenze.

- **Semplificazione dell'operatività.** Le console native delle piattaforme cloud sono strumenti pensati per specialisti. La maggior parte degli utenti ha bisogno di compiere poche operazioni ricorrenti — creare una macchina virtuale, aggiungere spazio disco, aprire una porta di rete — senza dover conoscere il modello concettuale della piattaforma sottostante.
- **Visione unificata dei consumi.** Chi eroga il servizio deve sapere, in ogni momento, quanta capacità è disponibile, quanta ne è impegnata e da chi. Chi il servizio lo utilizza deve sapere quanto del proprio piano ha consumato. Queste due informazioni oggi vivono in sistemi separati.
- **Indipendenza tecnologica.** Vincolarsi a una singola tecnologia o a un singolo fornitore è un rischio strategico, tanto più per soggetti che erogano servizi di interesse pubblico. La possibilità di aggiungere, sostituire o affiancare ambienti diversi senza riscrivere la piattaforma di gestione è un requisito, non un'opzione.
- **Sovranità del dato.** Dati, credenziali, metriche e automazioni devono restare nel perimetro dell'organizzazione e sotto giurisdizione nazionale ed europea. Questo vale anche — e in modo particolarmente stringente — per le componenti di intelligenza artificiale.

La CMP di Epic Edge è stata progettata a partire da queste quattro esigenze, e non come adattamento di un prodotto preesistente.

3. Che cos'è la Epic Edge CMP

La piattaforma si presenta come un portale web al quale l'utente accede con le proprie credenziali aziendali. Da quel momento in poi, tutto ciò che vede è determinato da chi è e da cosa ha effettivamente acquistato o gli è stato assegnato: la piattaforma non espone menu generici da filtrare, ma costruisce l'interfaccia intorno all'utente.

Caratteristica	Descrizione
Multi-provider	Gestisce ambienti cloud di natura diversa — piattaforme open source, virtualizzazione tradizionale, operatori pubblici — presentandoli tutti nello stesso modo all'utente.
White-label	La piattaforma è progettata perché l'operatore che la adotta la proponga ai propri clienti con il proprio marchio e la propria identità visiva; la personalizzazione white-label completa è in fase di rilascio. Epic Edge resta il fornitore della tecnologia e, in quella modalità, non compare all'utente finale.
Guidata dai dati	L'aggiunta di un nuovo ambiente o di un nuovo prodotto a catalogo non richiede un nuovo rilascio del portale: la piattaforma disegna ciò che le viene comunicato.
Assistita da AI	Due componenti integrate: EdgeForge AI per l'esecuzione delle operazioni in linguaggio naturale e CloudSentinel AI per la sorveglianza intelligente e l'ottimizzazione energetica, entrambe operanti nel perimetro dell'organizzazione.
Open source	L'intera catena tecnologica è basata su componenti open source, senza licenze proprietarie vincolanti e senza dipendenze da fornitori extraeuropei.
Installabile in loco	La piattaforma può essere installata nell'infrastruttura del cliente, mantenendo dati e credenziali sotto il suo esclusivo controllo.

3.1 Come funziona: l'architettura in breve

Tra l'interfaccia e gli ambienti gestiti opera un middleware multi-provider: è il componente che rende possibile presentare tecnologie diverse nello stesso modo. Traduce e normalizza le API di ciascun backend, espone alla piattaforma un catalogo uniforme di capability e affida l'esecuzione a un motore di automazione basato su Ansible e AWX. Gli assistenti AI non dialogano mai direttamente con l'infrastruttura: producono un piano, che il middleware esegue in modo tracciato. È questo strato di astrazione a permettere di aggiungere o sostituire un provider senza riscrivere gli assistenti né l'interfaccia.

Il modello propone. Il motore di automazione esegue.

La piattaforma è inoltre guidata dai dati: il portale non contiene schermate sviluppate ad hoc per ciascun provider o prodotto. Ogni ambiente descrive le proprie capability attraverso il middleware, e la CMP costruisce dinamicamente l'interfaccia, il catalogo e le operazioni disponibili. Introdurre un nuovo servizio significa quindi estendere il catalogo dati e gli adapter del middleware, non riscrivere l'interfaccia utente.

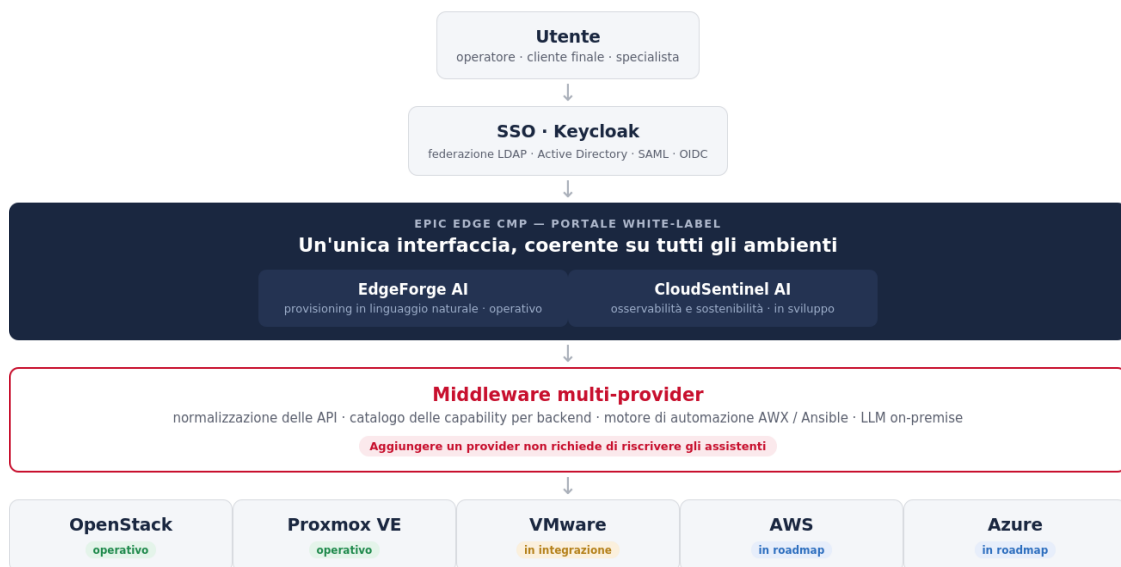


Figura 0 — Architettura logica: accesso federato, portale unico con gli assistenti AI, middleware multi-provider, backend eterogenei.

La CMP non replica le funzionalità delle piattaforme sottostanti: ne offre un modello operativo uniforme, lasciando disponibili le console native per tutte le funzioni specialistiche.

In che cosa è diversa da una CMP tradizionale. Esistono piattaforme di gestione multi-cloud consolidate, ma quasi tutte demandano l'automazione conversazionale e l'analisi a servizi esterni. Qui la differenza non sta nel numero di funzioni: sta nel fatto che l'intelligenza artificiale opera interamente on-premise — il modello non lascia il perimetro dell'organizzazione — su una catena interamente open source e reversibile. È la combinazione che rende l'automazione AI compatibile con requisiti di sovranità che una CMP commerciale, per come è costruita, non può soddisfare.

4. L'accesso alla piattaforma

L'ingresso avviene attraverso un sistema di autenticazione centralizzato, integrabile con i sistemi di identità già in uso presso l'organizzazione. L'utente utilizza un'unica credenziale per tutti gli ambienti, senza dover conoscere né gestire le utenze tecniche delle singole piattaforme sottostanti. L'autenticazione è basata su un provider di identità standard (Keycloak): la piattaforma può quindi federarsi con le directory aziendali esistenti — LDAP,

Active Directory — e con provider di terzi via SAML o OpenID Connect, riutilizzando le identità e le politiche di accesso già in uso presso l'organizzazione.

Al termine dell'autenticazione la piattaforma conosce già il profilo dell'utente, il suo ruolo e l'elenco dei servizi che gli competono. È questa informazione a determinare l'intera esperienza successiva.

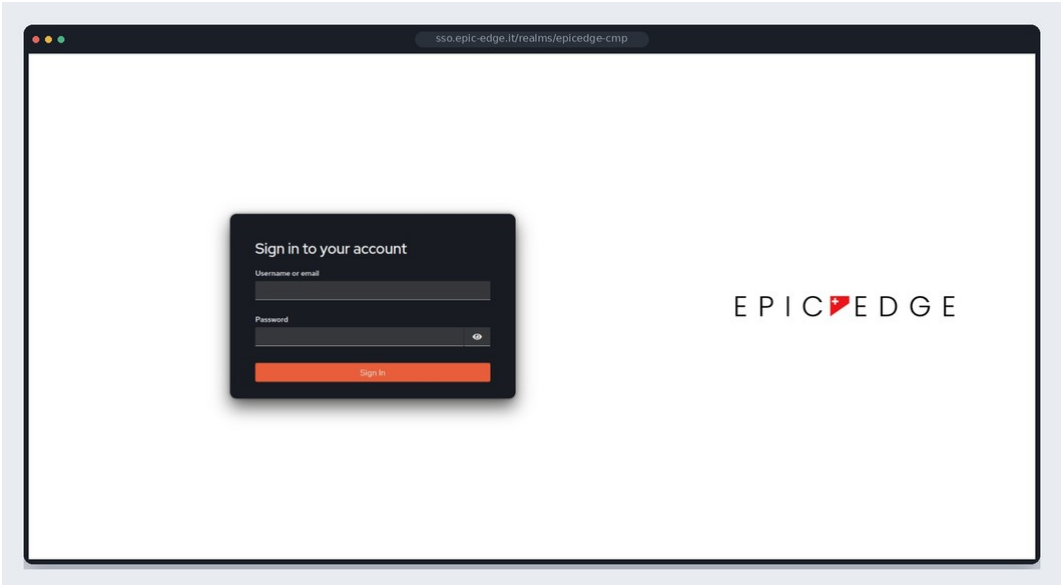


Figura 1 — Accesso unico via SSO: un'unica credenziale per tutti gli ambienti, federabile con le identità aziendali.

5. Due profili di utilizzo, un'unica piattaforma

La piattaforma distingue due profili, che accedono alla stessa interfaccia ma con finalità e visibilità differenti.

Profilo	Ruolo e ambito di visibilità
Amministratore	È l'operatore che eroga il servizio. Governa il catalogo dei prodotti offerti e accede a tutti gli ambienti in modalità di monitoraggio: vede la capacità complessiva di ciascun ambiente e i clienti che vi operano sopra, con i rispettivi consumi.
Cliente finale	È l'utilizzatore del servizio. Vede l'intera offerta disponibile, opera sugli ambienti che gli sono stati attivati e, per quelli non ancora attivi, viene indirizzato al percorso di acquisto definito dall'operatore.

5.1 La vista dell'amministratore

L'amministratore trova nella schermata iniziale due fasce distinte: gli ambienti tecnologici che offre ai propri clienti, ciascuno con l'indicazione di quanti clienti vi operano, e il catalogo dei prodotti commerciali che pubblica e gestisce in autonomia.

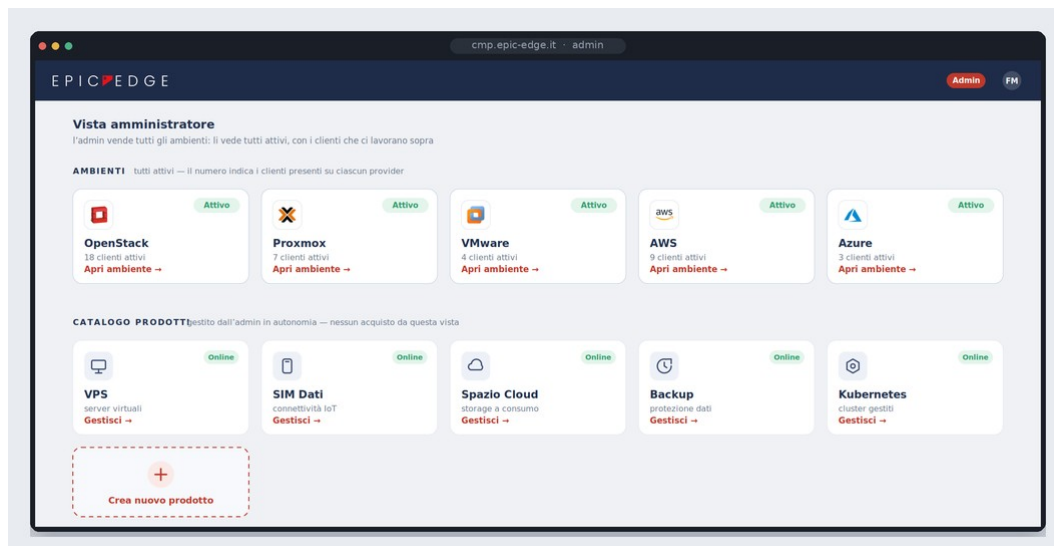


Figura 2 — Vista amministratore: gli ambienti offerti e il catalogo dei prodotti, governati da un unico punto.

Entrando in un ambiente, l'amministratore non vede risorse proprie ma la fotografia complessiva dell'ambiente: la capacità disponibile e quella impegnata per memoria, potenza di calcolo e spazio di archiviazione, il numero di istanze in esecuzione e l'elenco dei clienti attivi con i consumi di ciascuno. È la vista di chi deve pianificare la capacità e rendicontare il servizio.

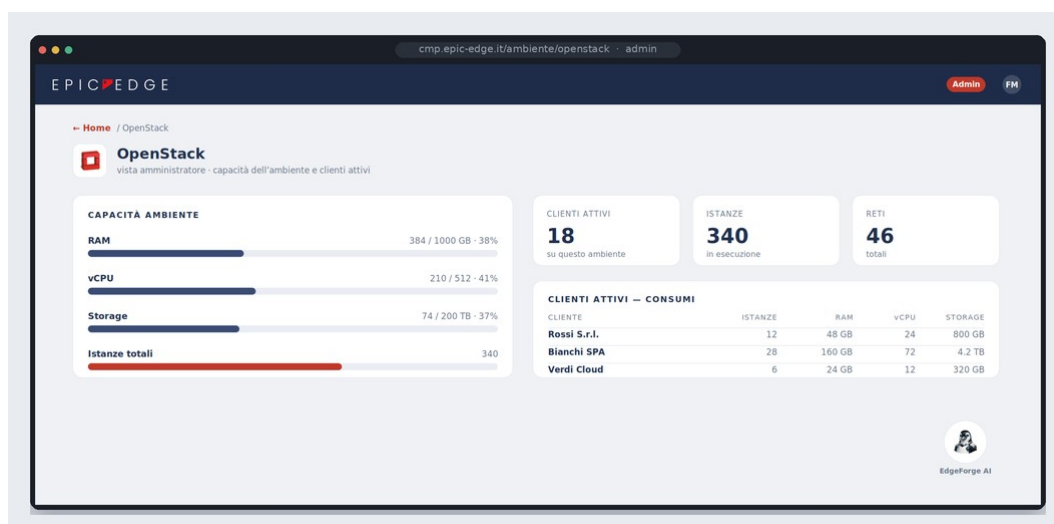


Figura 3 — Vista amministratore di un ambiente: capacità disponibile e consumi per singolo cliente, per pianificare e rendicontare.

La piattaforma espone inoltre, nello stesso portale, le dashboard di osservabilità dell'ambiente, costruite sugli strumenti standard già presenti — Prometheus e Grafana. Le metriche vengono raccolte valorizzando l'infrastruttura di osservabilità già presente nell'ambiente, senza introdurre agenti proprietari nelle macchine virtuali e senza

uplicare la raccolta. Accanto ai dati aggregati di capacità e consumo, l'amministratore dispone delle metriche di dettaglio per singola risorsa — CPU, memoria, rete e disco delle macchine virtuali — e di una lettura orientata al capacity planning: occupazione per ambiente, andamento nel tempo e margine residuo, come base per pianificare gli ampliamenti. La piattaforma è inoltre progettata per offrire una vista aggregata cross-provider che riunisce in un unico punto le metriche di ambienti tecnologicamente diversi, senza aprire una console di monitoraggio separata per ciascuno; questa funzionalità è attualmente in fase di implementazione nel portale unificato.

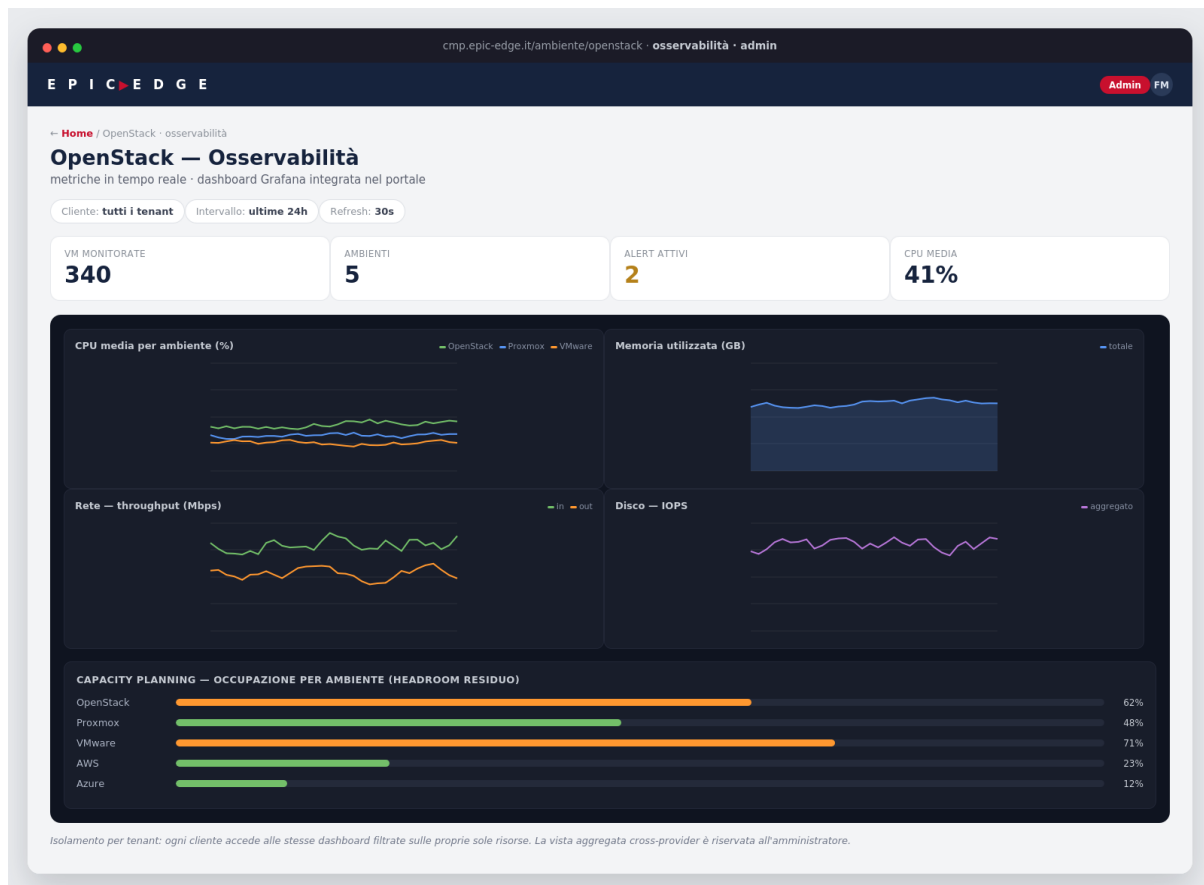


Figura 3-bis — Osservabilità integrata: dashboard Grafana con le metriche di dettaglio delle risorse e la lettura di capacity planning per ambiente.

Anche per l'amministratore è disponibile l'assistente EdgeForge AI, che in questo profilo non esegue operazioni di provisioning ma risponde a domande di monitoraggio: quali ambienti si stanno avvicinando alla saturazione, come si confrontano tra loro i diversi provider, quali risorse sono allocate a quale cliente. È in questa stessa area che sarà integrato CloudSentinel AI, descritto al capitolo 7.

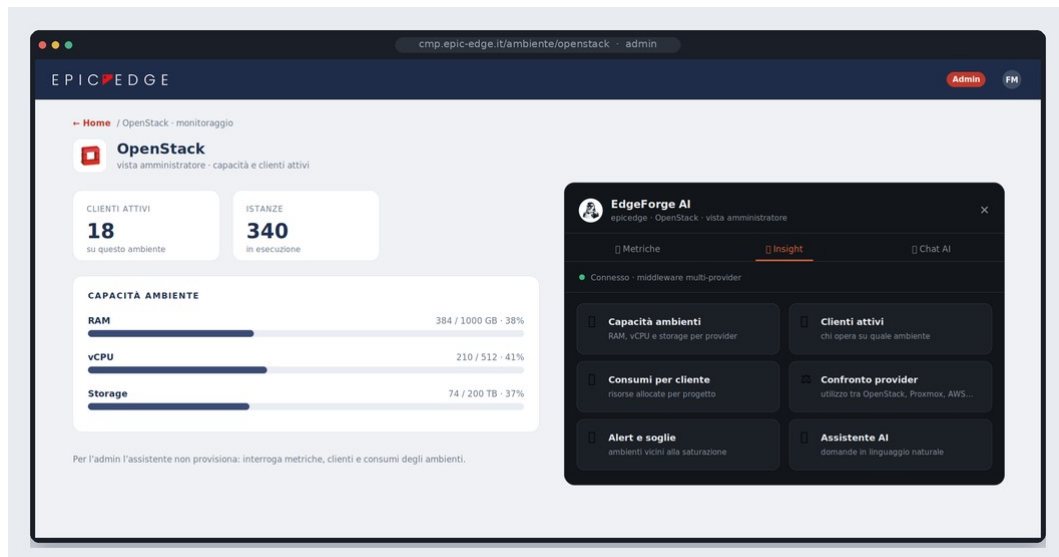


Figura 4 — Nella vista amministratore l'assistente AI risponde a domande di monitoraggio: non esegue provisioning.

5.2 La vista del cliente finale

Il cliente finale vede l'intera offerta dell'operatore. Gli ambienti che ha già attivato sono immediatamente accessibili; quelli non ancora attivi restano visibili e, se selezionati, conducono alla pagina di acquisto predisposta dall'operatore. La stessa logica vale per il catalogo dei prodotti. In questo modo l'interfaccia non è soltanto uno strumento operativo, ma anche un canale commerciale.

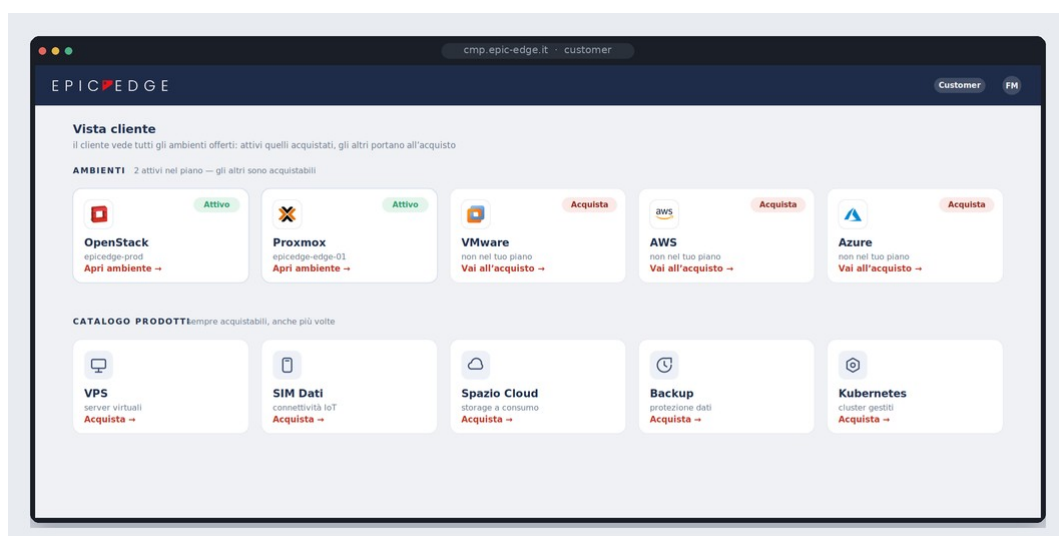


Figura 5 — Vista cliente: accesso agli ambienti attivi e al catalogo dei servizi acquistabili, nello stesso posto.

All'interno di un ambiente attivo il cliente trova soltanto ciò che gli appartiene: le proprie macchine, le proprie reti, i propri volumi e l'utilizzo del piano sottoscritto rispetto ai limiti concordati. La struttura della pagina è identica per tutti gli ambienti: cambia il contenuto, non il modo di leggerlo.

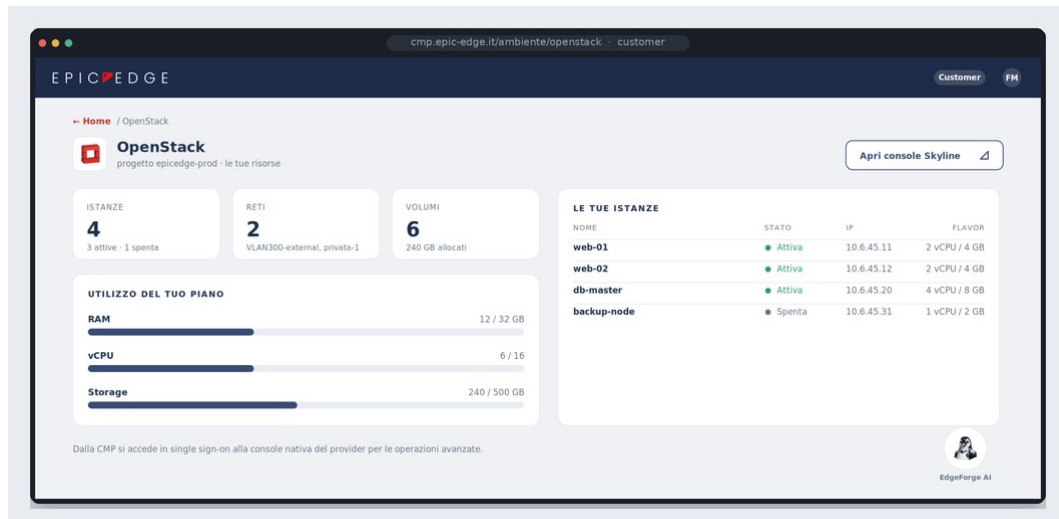


Figura 6 — Vista cliente di un ambiente: solo le proprie risorse e l'utilizzo del piano, con lettura identica per ogni provider.

Allo stesso modo, il cliente accede alle dashboard di osservabilità delle proprie risorse, con le metriche di dettaglio delle proprie istanze e dei propri servizi — CPU, memoria, rete e disco. È inoltre prevista una vista aggregata dei consumi complessivi sui diversi ambienti attivati, in fase di implementazione insieme al portale unificato. Le dashboard sono isolate per tenant: ogni cliente vede esclusivamente le proprie risorse, mai quelle di altri clienti sullo stesso ambiente.

Da questa stessa pagina il cliente può aprire l'assistente EdgeForge AI, che eredita automaticamente il contesto dell'ambiente in cui si trova e propone le automazioni disponibili raggruppate per aree funzionali, oltre alla possibilità di formulare richieste libere in linguaggio naturale.

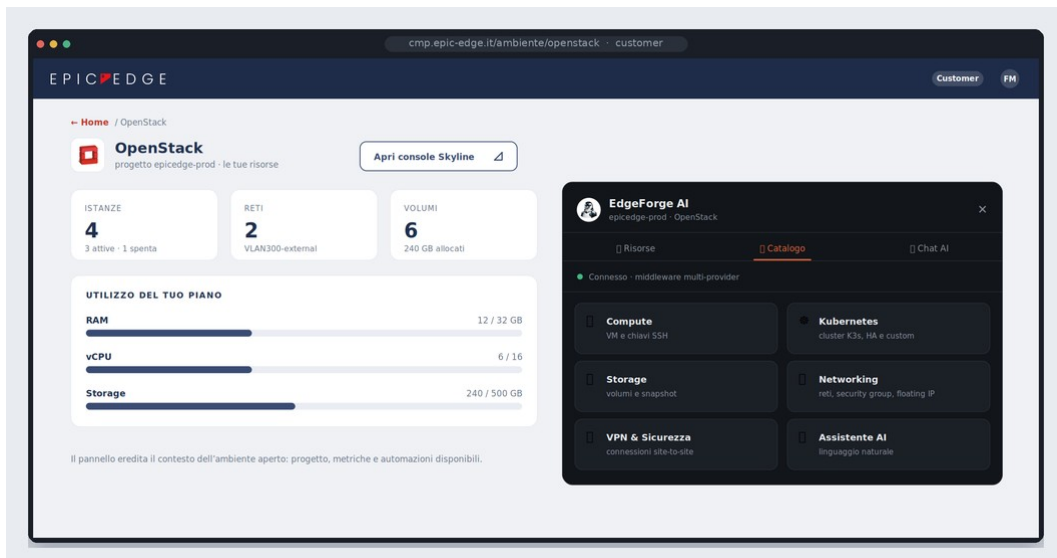


Figura 7 — L'assistente AI eredita il contesto dell'ambiente attivo e propone solo le automazioni realmente disponibili.

6. EdgeForge AI: l'assistente operativo

EdgeForge AI è la componente che distingue maggiormente la piattaforma. Non è un chatbot informativo affiancato al portale, ma un assistente operativo che esegue realmente le richieste dell'utente sull'infrastruttura, attraverso automazioni verificate e tracciate.

L'utente può descrivere ciò che desidera in linguaggio naturale — creare un ambiente di collaudo, aggiungere spazio a una macchina esistente, installare un applicativo su un server già attivo — e l'assistente traduce la richiesta in una sequenza di operazioni, la esegue e ne riporta l'avanzamento in tempo reale. In alternativa, per le operazioni ricorrenti, l'assistente propone percorsi guidati che accompagnano l'utente passo dopo passo. In ogni caso, EdgeForge AI propone sempre il piano operativo all'utente prima dell'esecuzione, consentendone la verifica e l'approvazione: nessuna risorsa viene creata senza un'approvazione esplicita.

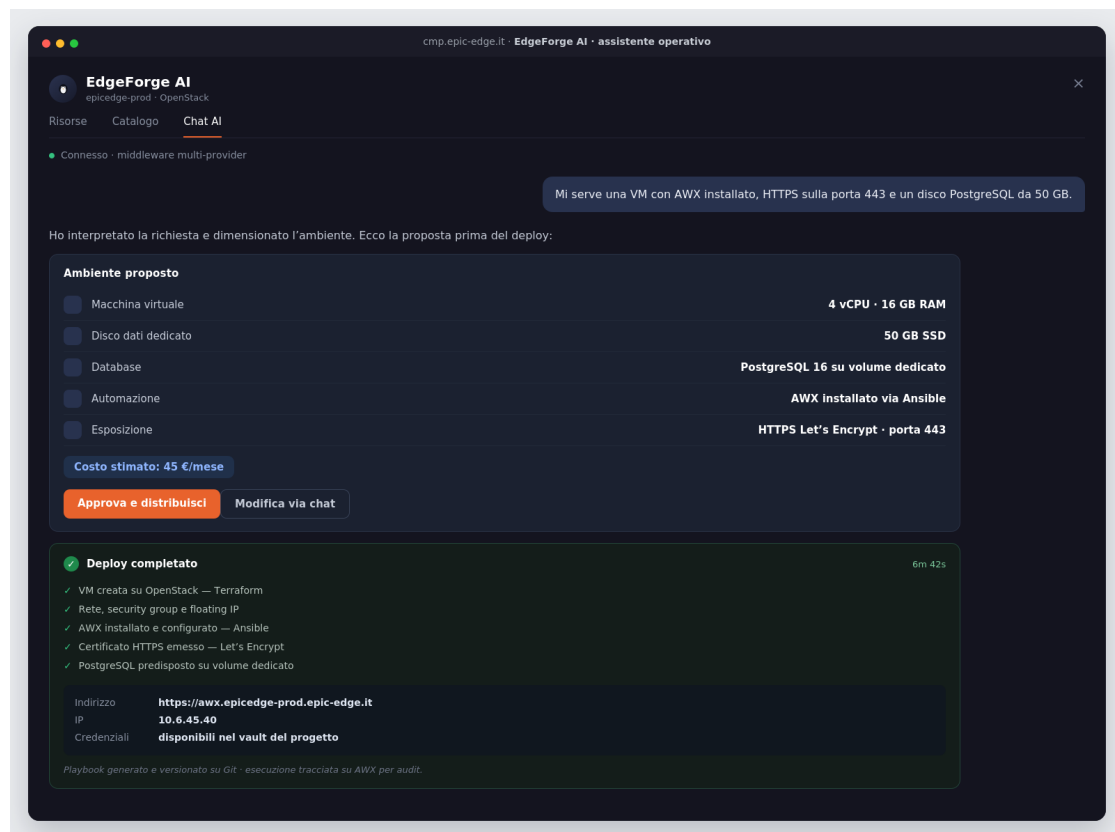
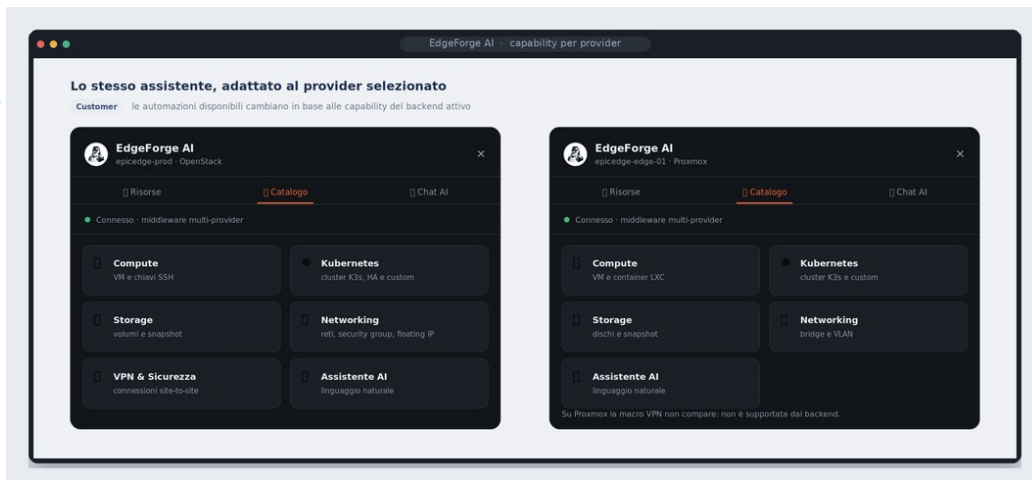


Figura 7-bis — EdgeForge AI in azione: dalla richiesta in linguaggio naturale alla proposta dell'ambiente, fino all'esecuzione automatizzata e tracciata del piano approvato.

- **Contestuale.** Sa in quale ambiente si trova l'utente e con quali autorizzazioni opera. Non è necessario specificarlo a ogni richiesta.
- **Multi-piattaforma.** Il provisioning automatico è già operativo su più tecnologie: oltre agli ambienti cloud open source, EdgeForge AI crea e gestisce risorse anche su piattaforme di virtualizzazione tradizionale, con la stessa interfaccia e lo stesso linguaggio. L'utente non deve sapere su quale tecnologia sta lavorando. Allo stato attuale il provisioning automatico è pienamente operativo su OpenStack e su Proxmox VE; l'estensione a VMware è in corso di integrazione.
- **Adattivo.** Le automazioni proposte cambiano in funzione delle reali capacità dell'ambiente selezionato. Se una funzionalità non è supportata dalla tecnologia sottostante, semplicemente non viene offerta: l'utente non incontra opzioni che non potrebbero funzionare.
- **Sovrano.** L'elaborazione avviene all'interno del perimetro dell'organizzazione. Nessun dato relativo all'infrastruttura, ai clienti o alle credenziali transita verso servizi di intelligenza artificiale esterni.
- **Tracciato.** Ogni operazione eseguita dall'assistente passa attraverso il sistema di automazione dell'organizzazione, che ne conserva la registrazione completa a fini di verifica e di audit.

La figura seguente mostra lo stesso assistente su due ambienti tecnologicamente diversi: in entrambi il provisioning automatico è pienamente funzionante, ma l'elenco delle automazioni si adatta a ciò che ciascuna piattaforma è realmente in grado di fare.

Figura 8 — Lo stesso



assistente AI si adatta alle capability del provider: su Proxmox la macro VPN non compare perché il backend non la supporta.

Perché questo aspetto è rilevante

L'adozione dell'intelligenza artificiale in ambito pubblico incontra spesso un ostacolo di natura non tecnica: l'impossibilità di far uscire dati sensibili dal perimetro dell'organizzazione. EdgeForge AI è progettato per operare interamente all'interno di quel perimetro, rendendo l'automazione conversazionale compatibile con i requisiti di riservatezza e di sovranità del dato.

L'estensione a VMware segue lo stesso principio: l'assistente presenta le risorse native dell'ambiente — datastore, distributed switch, template — e offre soltanto le automazioni effettivamente supportate dal backend, senza che l'utente debba conoscere la tecnologia sottostante.

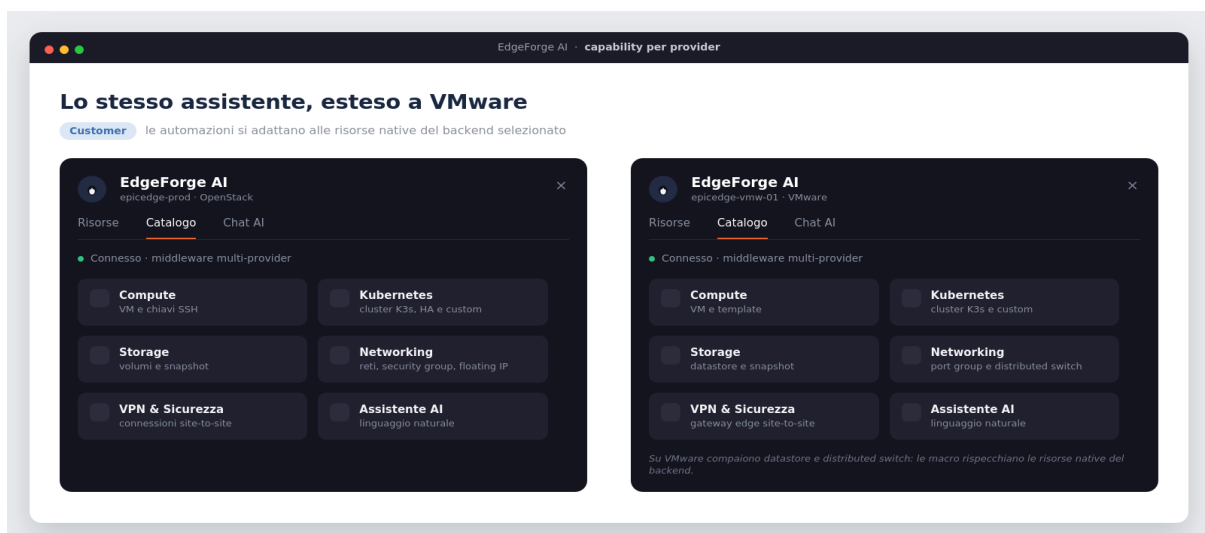


Figura 8-bis — L'assistente esteso a VMware: le automazioni rispecchiano le risorse native del backend (datastore, distributed switch, template).

7. CloudSentinel AI: sorveglianza intelligente e sostenibilità

Se EdgeForge AI si occupa di costruire e modificare l'infrastruttura, CloudSentinel AI si occupa di ciò che accade dopo: tenerla in salute, anticiparne i problemi e ridurre il consumo energetico. È la seconda componente di intelligenza artificiale della piattaforma, attualmente in fase di sviluppo, e sarà integrata nell'interfaccia dell'amministratore accanto alle attuali funzioni di monitoraggio.

CloudSentinel AI non sostituisce gli strumenti di monitoraggio già in uso presso l'organizzazione: li mette a sistema e li arricchisce con una capacità di analisi che oggi richiede competenze specialistiche rare e costose.

7.1 Le esigenze a cui risponde

Criticità	Impatto per l'organizzazione
Interruzioni non pianificate	Un fermo di servizio su infrastruttura critica ha un costo orario molto elevato e, nel settore pubblico, un impatto diretto sui cittadini. Gli impegni di continuità del servizio lasciano margini di indisponibilità dell'ordine di poche decine di minuti all'anno.
Tempi di ripristino lunghi	L'analisi della causa di un guasto complesso richiede oggi ore di lavoro di personale altamente specializzato, disponibile in numero limitato sul mercato.
Spreco energetico	L'utilizzo medio dei server nei centri di calcolo è tipicamente molto inferiore alla capacità installata: gran parte dell'energia assorbita non produce servizio.
Rendicontazione ambientale	Cresce l'obbligo, normativo e contrattuale, di documentare consumi ed emissioni dell'infrastruttura digitale, senza che esistano strumenti nativi per farlo sul cloud.

7.2 Le funzioni previste

CloudSentinel AI opererà lungo quattro direttrici, presentate all'amministratore all'interno della stessa interfaccia che già utilizza per il monitoraggio.

- **Rilevamento delle anomalie.** Impara il comportamento normale di ciascun ambiente e segnala gli scostamenti critici prima che diventino disservizio.
- **Previsione dei guasti.** Analizza l'andamento nel tempo degli indicatori dell'infrastruttura per anticipare i guasti con ore o giorni di preavviso, consentendo interventi programmati anziché reazioni d'emergenza.

- **Analisi automatica delle cause.** A incidente avvenuto, correla gli eventi delle diverse componenti e restituisce ricostruzione temporale, causa più probabile e azioni suggerite.
- **Rimedio assistito.** Propone ed esegue le azioni correttive con tre livelli di autonomia selezionabili dall'organizzazione: semplice segnalazione all'operatore, esecuzione previa approvazione, oppure esecuzione automatica limitata alle sole azioni a basso rischio.

7.3 Il modulo di sostenibilità energetica

A queste funzioni si affianca un modulo di efficienza energetica: consolida i carichi su un numero inferiore di server senza interruzione del servizio, tiene conto delle condizioni operative nel collocare i nuovi carichi, e produce automaticamente la rendicontazione di consumi ed emissioni per singolo servizio, in un formato pensato per essere coerente con i quadri europei di sostenibilità.

Il valore combinato delle due componenti

EdgeForge AI e CloudSentinel AI coprono insieme l'intero ciclo di vita dell'infrastruttura: descrivere e realizzare, poi sorvegliare e ottimizzare. Per l'organizzazione questo significa ridurre contemporaneamente il tempo necessario a erogare un servizio, il tempo necessario a ripristinarlo in caso di guasto e l'energia impiegata per mantenerlo attivo — tre grandezze che oggi vengono governate con strumenti separati e competenze diverse.

8. Continuità con le console tecniche

La piattaforma non sostituisce gli strumenti specialistici, li affianca. Per le operazioni avanzate l'utente accede, senza nuove autenticazioni, alla console nativa dell'ambiente in cui sta lavorando, dove ritrova l'intero insieme delle funzionalità della piattaforma sottostante.

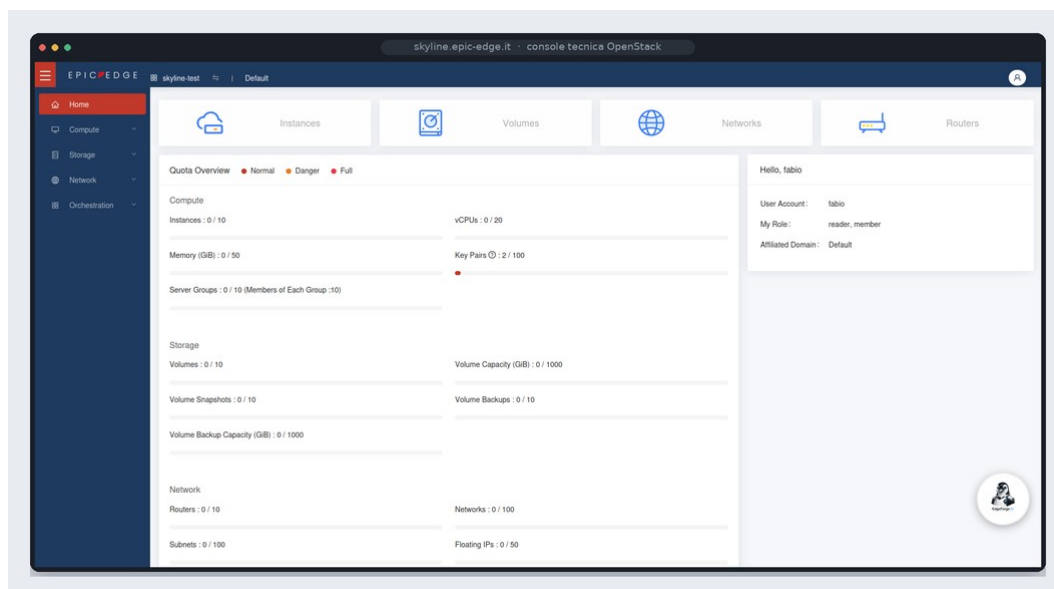


Figura 9 — Dalla CMP si accede in single sign-on alla console tecnica nativa del provider per le operazioni avanzate.

Anche in questo contesto l'assistente EdgeForge AI resta disponibile e mantiene le stesse funzionalità: l'utente specialista può continuare a lavorare con gli strumenti che conosce e, quando conviene, delegare all'assistente le operazioni più ripetitive. Le immagini seguenti documentano una funzionalità già operativa negli ambienti Epic Edge.

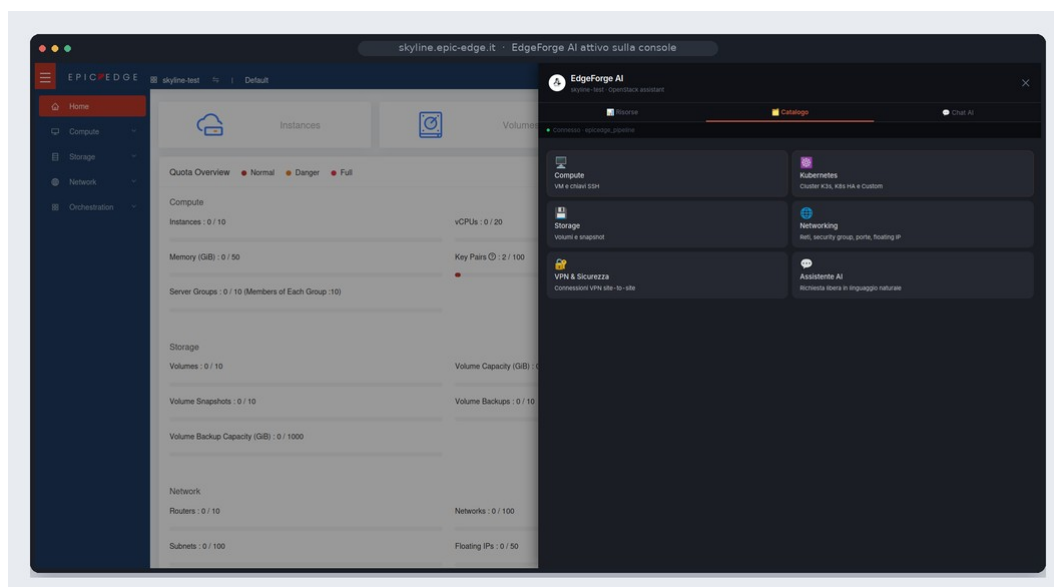


Figura 10 — L'assistente AI resta disponibile anche dentro la console tecnica: funzione già operativa negli ambienti Epic Edge.

8.1 Integrazione con l'ecosistema dell'organizzazione

La piattaforma non vive isolata: è progettata per inserirsi negli strumenti già presenti nell'organizzazione, riutilizzandoli anziché sostituirli. I principali punti di integrazione sono:

Identità. Federazione con i provider di identità esistenti — LDAP, Active Directory, SAML, OpenID Connect — tramite il sistema SSO.

Automazione. Esecuzione attraverso Ansible e AWX, con i playbook versionati su repository Git dell'organizzazione.

Monitoraggio. Raccordo con le metriche e gli strumenti di osservabilità già in uso, a partire da Prometheus.

Interoperabilità applicativa. API REST e notifiche via webhook per il collegamento con sistemi di ticketing e portali interni.

Alcuni raccordi sono già operativi (identità, automazione, metriche); altri, come i connettori verso specifici sistemi di ticketing, sono previsti nelle evoluzioni successive.

9. Il valore in un contesto di cloud nazionale

La piattaforma risponde in modo diretto ad alcuni dei principi che orientano le politiche nazionali ed europee sul cloud per la pubblica amministrazione.

Principio	Come la piattaforma vi risponde
Sovranità del dato	Installazione presso l'infrastruttura del cliente, giurisdizione italiana ed europea, nessuna dipendenza da servizi extraeuropei, incluse le componenti di intelligenza artificiale.
Reversibilità	Tecnologie open source e formati standard. L'organizzazione può cambiare fornitore o tecnologia sottostante senza dover riprogettare la propria piattaforma di gestione.
Assenza di vincoli	Nessuna licenza proprietaria vincolante e nessun costo di uscita. La piattaforma è progettata per convivere con ambienti eterogenei, non per indurre concentrazione su una sola tecnologia.
Continuità del servizio	CloudSentinel AI (in fase di sviluppo) è concepito per far evolvere la gestione dei guasti da reattiva a predittiva, con l'obiettivo di sostenere il rispetto degli impegni di disponibilità del servizio.
Tracciabilità	Ogni operazione eseguita tramite la piattaforma, sia manuale sia assistita, è registrata dal sistema di automazione e disponibile per verifiche e audit.
Accessibilità operativa	Riduzione della barriera di competenza necessaria per operare sul cloud, con effetti diretti sui tempi di erogazione e sui costi di formazione del personale.

Principio	Come la piattaforma vi risponde
Trasparenza dei consumi	Visione unificata della capacità disponibile e di quella impegnata, per singolo ambiente e per singolo utilizzatore, come base per la pianificazione e la rendicontazione.
Sostenibilità ambientale	Il modulo energetico di CloudSentinel AI (in fase di sviluppo) è progettato per ridurre il consumo dell'infrastruttura e per produrre la rendicontazione ambientale richiesta dai quadri normativi europei.

10. Stato del progetto

Il progetto si trova in una fase avanzata di sviluppo, con componenti già in esercizio e altre in corso di realizzazione.

- **Già operativo.** L'assistente EdgeForge AI è funzionante e in uso all'interno degli ambienti Epic Edge, integrato nella console tecnica e collegato al sistema di automazione. Il provisioning automatico è già attivo su piattaforme cloud open source (OpenStack) e su virtualizzazione tradizionale (Proxmox VE). Le figure 9 e 10 documentano la situazione reale.
- **In realizzazione.** Il portale unificato multi-provider, di cui le figure da 1 a 8 rappresentano la progettazione dell'esperienza utente approvata e in corso di implementazione. L'estensione del provisioning automatico a ulteriori ambienti, VMware in primo luogo, è in corso di integrazione.
- **In sviluppo.** CloudSentinel AI, la cui architettura è definita e per il quale è già avviata la raccolta dei dati di esercizio necessari all'addestramento dei modelli. L'integrazione avverrà nell'interfaccia dell'amministratore, accanto alle funzioni di monitoraggio già presenti.
- **Evoluzioni previste.** Estensione del numero di tecnologie supportate, personalizzazione grafica avanzata per singolo operatore e ampliamento delle funzionalità di analisi predittiva dei consumi.

Nota sulle immagini

Le figure 9 e 10 sono immagini di un ambiente reale in esercizio. Le figure da 1 a 8 rappresentano la progettazione dell'interfaccia del portale unificato, attualmente in fase di realizzazione. CloudSentinel AI, descritto al capitolo 7, è in fase di sviluppo e non è ancora rappresentato nelle immagini. Il branding mostrato nelle figure è quello dell'istanza dimostrativa Epic Edge; in modalità white-label è sostituibile con l'identità dell'operatore. La figura 8-bis illustra l'estensione dell'assistente all'ambiente VMware ed è anch'essa parte della progettazione in corso.

10.1 Quadro sinottico dello stato

Funzione	Stato
Accesso SSO (Keycloak, LDAP/AD/SAML/OIDC)	● Disponibile
OpenStack — provisioning EdgeForge AI	● Disponibile
Proxmox VE — provisioning EdgeForge AI	● Disponibile
EdgeForge AI — assistente operativo	● Disponibile
Portale unificato multi-provider	● In implementazione
VMware — provisioning	● In integrazione
AWS — gestione dell'ambiente	● In roadmap
Azure — gestione dell'ambiente	● In roadmap
White-label	● In rilascio
CloudSentinel AI — osservabilità e sostenibilità	● In sviluppo
Osservabilità per ambiente (dashboard Grafana)	● Disponibile
Metriche aggregate cross-provider + capacity planning	● In implementazione

Legenda: ● verde = disponibile e in uso · ● ambra = in implementazione o rilascio · ● blu = in sviluppo.

11. Perché questa piattaforma è strategica per un Cloud Service Provider nazionale

Al netto delle singole funzioni, il valore della piattaforma per un operatore nazionale si riassume in cinque punti concreti:

Un unico front-end. OpenStack, Proxmox, VMware e le ulteriori piattaforme integrate nel tempo attraverso il middleware multi-provider, governate con la stessa esperienza d'uso, senza moltiplicare le console e la formazione.

Provisioning più rapido. L'automazione in linguaggio naturale porta la messa in esercizio di un ambiente standard da giorni di lavoro manuale a un deploy automatizzato di pochi minuti.

Nessuna dipendenza da AI esterne. Il modello resta nel perimetro dell'organizzazione: l'automazione intelligente non impone di esporre dati e infrastruttura a servizi extraeuropei.

White-label per l'operatore nazionale. Il servizio è proposto al cliente finale con il marchio e l'identità dell'operatore, non con quelli del fornitore della tecnologia.

Backend che evolve in modo indipendente. Aggiungere, sostituire o dismettere una tecnologia sottostante non modifica l'esperienza dell'utente finale.

L'ecosistema Epic Edge



Ogni prodotto Epic Edge fa parte di SentinelForge AI: moduli distinti sullo stesso substrato — LLM on-premise, esecuzione tracciata, policy di autonomia, audit, middleware multi-provider.

12. Epic Edge

Epic Edge S.r.l. è una società italiana specializzata nella progettazione, realizzazione e gestione di infrastrutture cloud basate su tecnologie open source. Opera su ambienti di produzione critici per clienti enterprise e operatori di servizi in Italia e in Europa, con competenze specifiche su piattaforme cloud open source, storage distribuito, orchestrazione di container e servizi gestiti con livelli di servizio garantiti.

Il posizionamento dell'azienda è centrato sulla sovranità del dato e sull'indipendenza tecnologica: infrastrutture interamente open source, dati in Italia e in Unione Europea, nessun vincolo contrattuale o tecnico che ostacoli la reversibilità.



Epic Edge S.r.l.

www.epic-edge.it · info@epic-edge.it