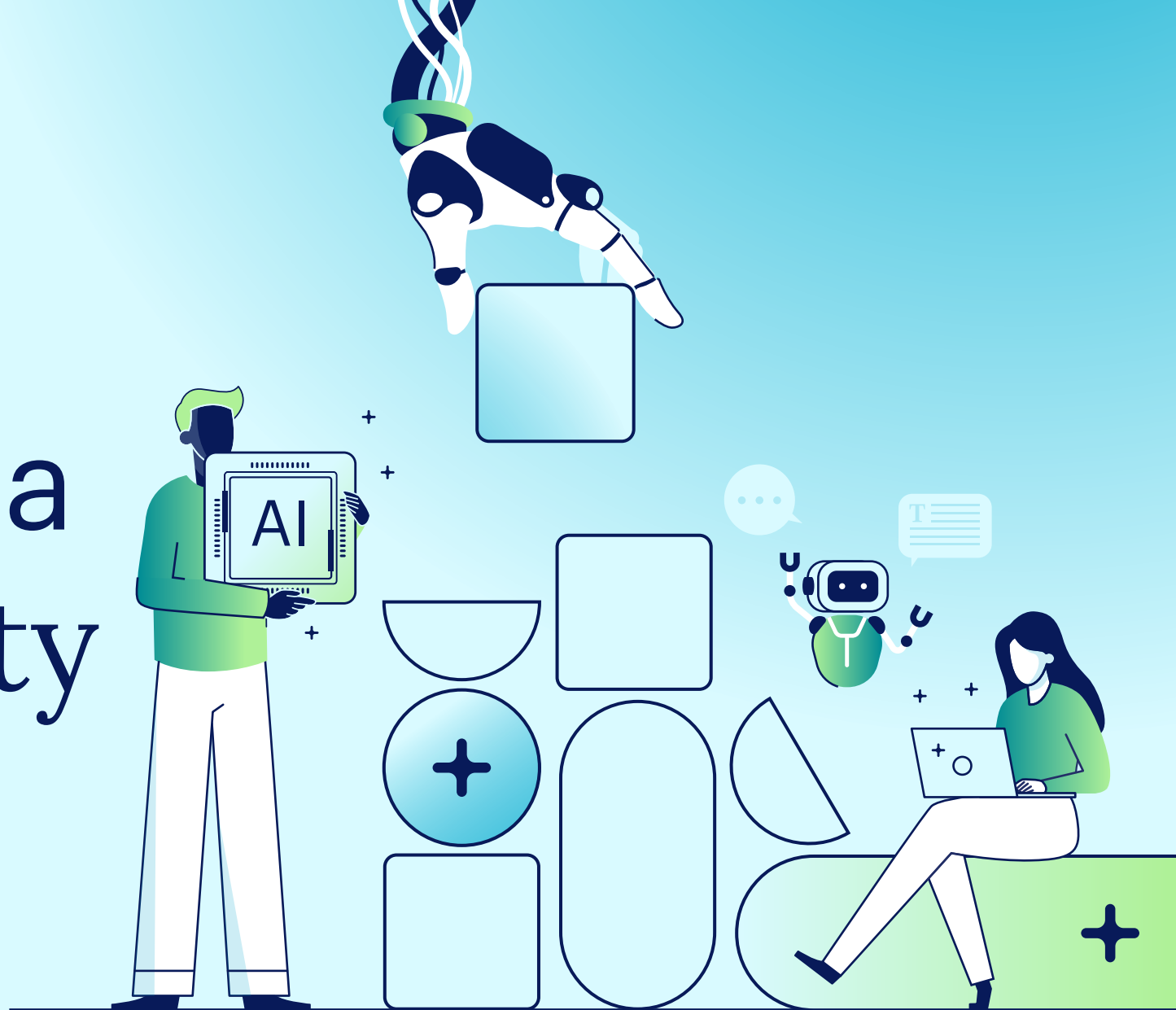




II

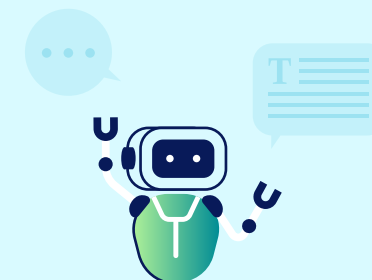
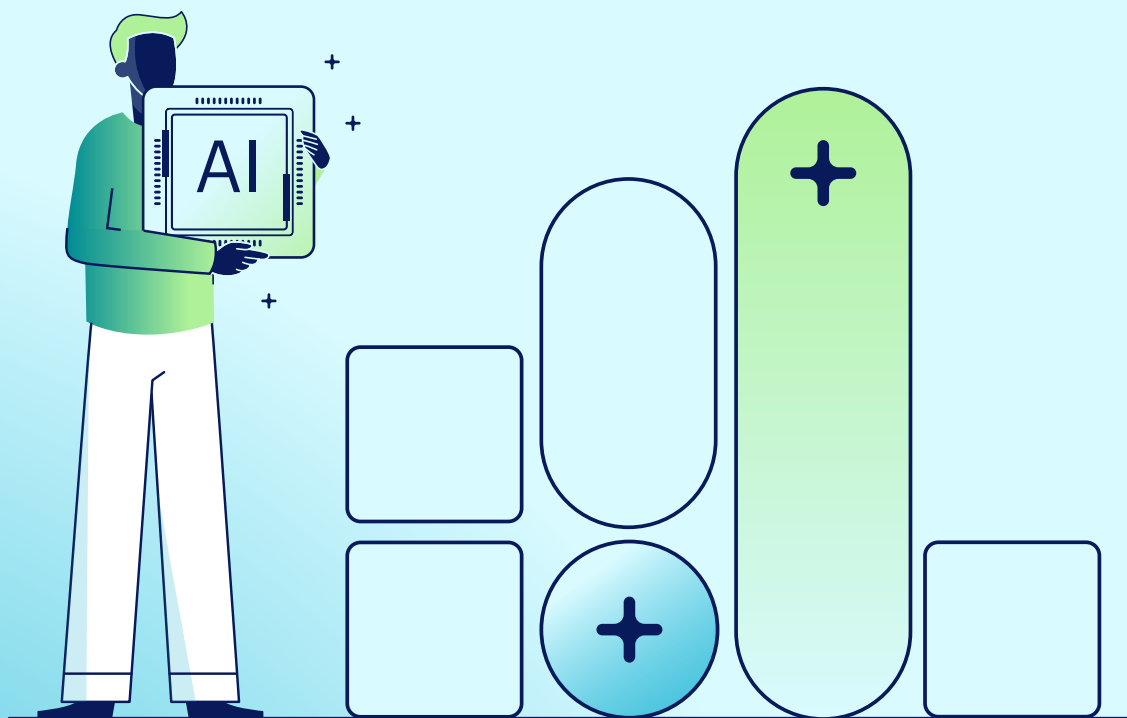
Manuale della AI Security

+ Una guida pratica per proteggere
l'AI end-to-end, ovunque



II

Manuale della AI Security



Indice

Introduzione	3
Sfide di sicurezza nell'AI	4
Fondamenti di sicurezza dell'AI	5
Navigare nella sicurezza dell'AI	6
Il futuro della sicurezza dell'AI	12
Conclusioni	13
A proposito di Netskope	14

Introduzione

Le tecnologie di intelligenza artificiale (AI) si sono rapidamente affermate come strumenti utili e importanti per molte aziende. Con nuove capacità e casi d'uso in continuo aumento, adesso l'AI è un componente fondamentale della maggior parte degli stack tecnologici aziendali.

La rapida ascesa dell'AI è stata segnata anche da alti livelli di investimenti. Secondo gli analisti di IDC, si stima che il mercato della spesa IT mondiale per l'AI sfiorerà i 750 miliardi di dollari entro il 2028, con una spesa specifica per l'AI generativa di poco superiore a 300 miliardi di dollari¹.

+ Si stima che il mercato della spesa IT mondiale per l'AI sfiorerà i 750 miliardi di dollari entro il 2028, con una spesa specifica per l'AI generativa di poco superiore a 300 miliardi di dollari.

Per i professionisti della sicurezza, i potenziali rischi delle applicazioni AI nel loro ambiente sono evidenti e in crescita. I ricercatori di Netskope Threat Labs hanno scoperto che il codice sorgente rappresenta quasi la metà (48%) di tutte le violazioni relative alla prevenzione della perdita di dati tramite genAI, con i dati regolamentati che rappresentano il 23% e la

proprietà intellettuale il 17%². Dal punto di vista della sicurezza, ciò solleva domande sui dati che i dipendenti inseriscono in questi sistemi e sui controlli in atto per gestirli.

Le sfide di sicurezza sono destinate a intensificarsi man mano che la tecnologia AI aziendale si evolve. I sistemi di AI agentica, ad esempio, possono operare in modo autonomo per raggiungere obiettivi specifici o eseguire compiti definiti senza richiedere un intervento umano costante. Gli analisti del settore di Gartner prevedono che, entro il 2028, il 25% delle violazioni aziendali sarà legato all'abuso di agenti AI³.

Considerata la rapida evoluzione dei rischi, non sorprende che i professionisti della sicurezza cerchino aiuto per orientarsi nel nuovo panorama. In questo eBook descriviamo le principali preoccupazioni di sicurezza che le aziende affrontano oggi e le soluzioni che Netskope può fornire.

+ Gartner prevede che, entro il 2028, il 25% delle violazioni aziendali sarà legato all'abuso di agenti AI.



¹ IDC Market Forecast, Worldwide Artificial Intelligence IT Spending Forecast, 2024–2028, Rick Villars et al., Ottobre 2024, n. doc. US52635424.

² Netskope Cloud and Threat Report, 2025 <https://www.netskope.com/netskope-threat-labs/cloud-threat-report/cloud-and-threat-report-2025>

³ Principali previsioni di Gartner per il 2025.

Sfide di sicurezza nell'AI

I tre principali problemi che affrontano oggi i team di sicurezza

1 Espansione della superficie del rischio

Con l'evolversi dell'uso dell'AI da strumento di AI generativa puro (come ChatGPT) a capacità di AI integrata nelle applicazioni aziendali e nelle applicazioni AI costruite privatamente, la superficie di attacco continua a espandersi. Ogni fase introduce nuovi rischi:

- Gli strumenti genAI pubblici sono largamente usati dai dipendenti e introducono rischi di esposizione involontaria di dati sensibili.
- Le funzioni AI integrate nelle app SaaS o cloud esistenti possono aprire percorsi nascosti per la fuoriuscita o la manipolazione dei dati.
- Gli LLM ospitati privatamente e le app AI personalizzate introducono nuovi vettori, come controlli di accesso configurati male o vulnerabilità nei prompt dei modelli e nelle pipeline di dati.

2 Esposizione ed estrazione dei dati sensibili

Il rischio più immediato nell'adottare l'AI è la perdita di dati, sia accidentale che malevola:

- L'esposizione involontaria si ha quando i dipendenti inseriscono dati sensibili (ad es., PII, segreti commerciali, dati regolamentati) in modelli pubblici senza rendersi conto delle conseguenze.
- Insider malevoli o attaccanti esterni possono sfruttare strumenti AI per estrarre dati o abusare dei canali di output del modello.
- C'è anche un rischio di addestramento: Addestrare modelli usando dati curati male può portare a rivelare informazioni riservate.

3 Governance responsabile dell'AI

Man mano che i sistemi AI si espandono, sollevano importanti questioni sulla conformità e sull'etica che si intersecano con la sicurezza:

- I modelli AI possono inavvertitamente codificare e propagare pregiudizi, portando a un'attenzione normativa e a danni d'immagine.
- Una gestione inadeguata dei dati dei dipendenti o dei clienti usati nei flussi di lavoro dell'AI potrebbe violare il GDPR, l'HIPAA o altre leggi sulla privacy dei dati.
- L'uso dell'AI al posto del processo decisionale umano, soprattutto in aree ad alto rischio (ad es., assunzioni, sicurezza, finanza), introduce dilemmi etici e carenze di responsabilità.

Fondamenti di sicurezza dell'AI

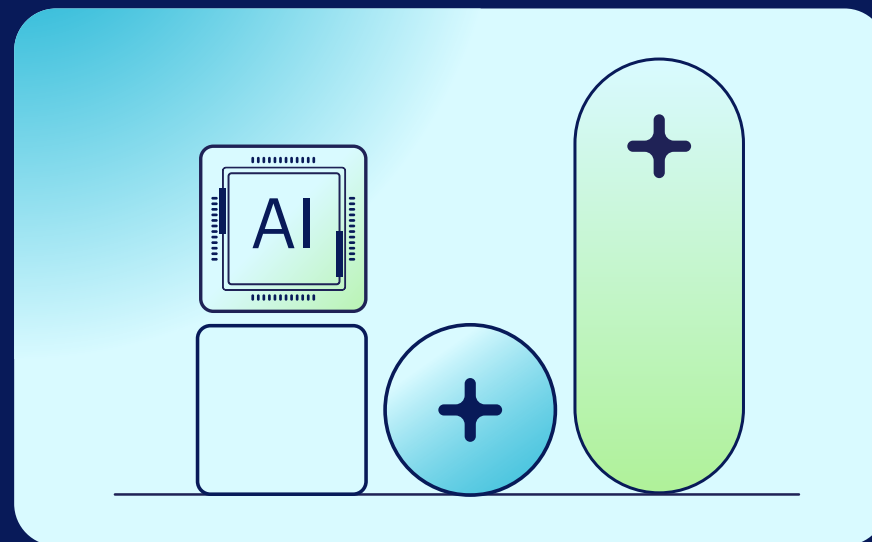
L'imperativo della zero trust

La sicurezza dell'AI si basa su un approccio di zero trust molto simile alla sicurezza SaaS, ma con sfide uniche che derivano dal modo in cui i modelli di AI elaborano gli input e generano output.

Sia la sicurezza dell'AI che quella del SaaS richiedono controlli di accesso rigorosi, un monitoraggio continuo e una protezione dei dati robusta per mitigare i rischi. Tuttavia, mentre la sicurezza SaaS si concentra principalmente sulla protezione delle applicazioni e delle interazioni degli utenti, la sicurezza dell'AI deve anche tenere conto dell'integrità dei dati di addestramento, dell'accesso ai modelli e della potenziale manipolazione avversaria. Questo rende essenziale applicare politiche di sicurezza consapevoli del contesto e rilevare minacce in tempo reale per prevenire la perdita di dati, l'accesso non autorizzato e lo sfruttamento dei modelli AI.

Un forte framework zero trust per la sicurezza dell'AI garantisce che ogni richiesta venga verificata, ogni flusso dati monitorato e l'accesso fornito sulla base di valutazioni del rischio dinamiche piuttosto che permessi statici. Questo approccio richiede una visibilità granulare nel movimento dei dati e controlli di sicurezza adattivi che si regolano in base al contesto in tempo reale.

Con i principi zero trust in atto, le aziende possono adottare e scalare in modo sicuro tecnologie guidate dall'AI senza compromettere la sicurezza o la conformità.



Consiglio del professionista

La sicurezza dell'AI si basa su un approccio zero trust molto simile alla sicurezza SaaS, ma con sfide uniche che derivano dal modo in cui i modelli di AI elaborano gli input e generano output.

Navigare nella sicurezza dell'AI

Le sei principali sfide e soluzioni



Sfida n. 1:

Mancanza di visibilità

Man mano che gli strumenti AI si integrano nei flussi di lavoro giornalieri, le aziende si trovano ad affrontare una sfida fondamentale in materia di sicurezza: non possono proteggere ciò che non possono vedere.

I dipendenti accedono ad applicazioni autorizzate e non autorizzate con credenziali aziendali e personali, sfumando i confini tra uso approvato e non approvato. Questa espansione incontrollata aumenta il rischio di perdita dei dati, perdita di proprietà intellettuale e violazioni della conformità, soprattutto quando informazioni sensibili vengono inserite in servizi AI non gestiti o shadow AI.

La maggior parte delle aziende non ha la visibilità granulare necessaria per differenziare tra l'uso rischioso e quello legittimo dell'AI. Gli strumenti tradizionali non riescono a individuare le interazioni specifiche dei modelli AI, a distinguere tra account personali e aziendali e a fornire informazioni in tempo reale a livello di utente, app o attività. Senza una visibilità approfondita su come e dove viene usata l'AI, i team di sicurezza rimangono ciechi di fronte ai potenziali punti di esposizione.



Come risolve il problema Netskope

Più le aziende adottano strumenti AI, più diventa fondamentale mantenere visibilità e controllo sul loro impiego. Netskope offre una soluzione completa per monitorare le applicazioni AI gestite e non gestite (shadow), fornendo ai team di sicurezza le informazioni necessarie per assicurare un'adeguata supervisione.

Le principali capacità includono:

- **Consapevolezza avanzata delle istanze:** Distinguere tra istanze personali e aziendali delle applicazioni AI come ChatGPT, Gemini e Copilot.
- **Dashboard AI:** Ottenere approfondimenti dettagliati su tendenze d'uso dell'AI, principali applicazioni, frequenza di accesso e azioni degli utenti a livello granulare, come accessi, post, upload e download.
- **Analisi del comportamento degli utenti e delle entità (UEBA, User and Entity Behavior Analytics):** Rilevare anomalie e comportamenti a rischio attraverso l'apprendimento automatico per individuare minacce come l'estrazione dei dati, i rischi interni e le violazioni delle politiche.
- **Visibilità granulare:** Monitorare e gestire il "cosa" e il "come" dell'AI dentro l'azienda, assicurando una protezione e una conformità solide.

Questa visibilità olistica permette ai team di sicurezza di intervenire in fretta e mitigare i rischi legati all'AI in tutta l'azienda.



Sfida n. 2:

Capire il rischio delle applicazioni AI

Con la rapida evoluzione delle capacità AI, anche il panorama dei rischi cambia. Quella che una volta era una semplice applicazione SaaS adesso può introdurre silenziosamente funzioni AI integrate come la generazione di testi, risposte intelligenti e copiloti AI, senza avvisare gli utenti o i team di sicurezza. Questa crescente tendenza rende sempre più difficile capire quali applicazioni usano l'AI, come la usano e quali rischi introducono nell'azienda.

I team di sicurezza hanno bisogno della capacità di valutare in modo dinamico il rischio in base a come sono integrate le funzioni dell'AI, se trattengono o addestrano dati aziendali e come si allineano ai requisiti di conformità. Senza questo livello di comprensione, le aziende rischiano di essere esposte a perdite di dati, furti di proprietà intellettuale, violazioni normative e anche manipolazioni dei modelli AI. Con l'espandersi dell'AI nel SaaS, capire il rischio delle applicazioni non è solo una buona pratica, ma una necessità per qualsiasi azienda che vuole adottare l'AI in modo sicuro.



Come risolve il problema Netskope

Netskope affronta la crescente complessità del rischio delle applicazioni AI con il suo Cloud Confidence Index (CCI), che fornisce informazioni in tempo reale, continuamente aggiornate, su oltre 83.000 applicazioni cloud e SaaS. Con valutazioni del rischio dinamiche e consapevoli dell'AI, CCI aiuta i team di sicurezza a rimanere un passo avanti rispetto ai rischi potenziali e ad assicurare la conformità.

Le principali capacità includono:

- **Valutazione del rischio in tempo reale, consapevole dell'AI:** Individuare le applicazioni con capacità AI integrata e capire i rischi associati a queste funzioni.
- **Approfondimenti sulla gestione dei dati aziendali:** Valutare come le applicazioni gestiscono i dati aziendali, inclusi la conservazione, l'addestramento dei modelli e la condivisione con terze parti.
- **Tracciamento della conformità:** Rimanere allineati ai requisiti normativi come GDPR, SOC 2 e ISO 27001.
- **Profili di rischio adattivi:** Quando il profilo di rischio di un'applicazione cambia a causa di una violazione, di termini aggiornati o di nuove funzioni AI, CCI si aggiorna in tempo reale, fornendo visibilità immediata su potenziali problemi di sicurezza.

Con CCI, i team di sicurezza possono navigare con fiducia nelle complessità dei rischi delle applicazioni AI e garantire che la loro azienda rimanga sicura e conforme.



Sfida n. 3: Integrità del modello AI

Più le aziende usano gli strumenti di AI generativa (sia modelli personalizzati che applicazioni aziendali come Microsoft Copilot), più diventa critico assicurare l'integrità dei dati per addestrare i modelli. Questi sistemi AI sono spesso addestrati su vasti set di dati, che possono includere documenti aziendali sensibili, e-mail, presentazioni, fogli di calcolo e informazioni commerciali riservate.

Se nei set dei dati di addestramento vengono incorporati dei dati sensibili o riservati, ciò può portare a un'esposizione non solo attraverso gli output del modello ma anche i prompt avversari, a perdita di dati e potenziali violazioni di conformità. Con l'espandersi della genAI in vari dipartimenti, diventa sempre più difficile per i team di sicurezza controllare come vengono reperiti, convalidati e protetti i dati di addestramento.

+ Microsoft Copilot può essere addestrato sui contenuti nella suite Office di un utente, dai documenti Word ai fogli Excel. Se in queste posizioni sono memorizzati dati riservati o sensibili e se i controlli di accesso non sono configurati bene, c'è il potenziale rischio che Copilot, nelle sue risposte, riveli strategie aziendali sensibili, dettagli finanziari o informazioni sui clienti.



Come risolve il problema Netskope

Netskope One DSPM (*Data Security Posture Management*, Gestione della postura di sicurezza dei dati) consente alle aziende di monitorare e proteggere i dati sensibili in ambienti cloud e repository di dati. Rilevando e classificando dati critici, come registri finanziari, PII e proprietà intellettuale, Netskope garantisce che queste informazioni non vengano usate per addestrare modelli AI senza la dovuta autorizzazione.

Tra le capacità principali:

- **Monitoraggio continuo degli ambienti cloud:** Rilevare e classificare i dati sensibili in tempo reale, garantendo che venga impedito l'uso non autorizzato nell'addestrare i modelli AI.
- **Visibilità nell'accesso e nella condivisione dei dati:** Ottenere informazioni in tempo reale sulle modalità di accesso e condivisione dei dati nel cloud, consentendo un'azione correttiva immediata quando necessario.
- **Conformità e prevenzione della perdita di dati:** Proteggere i dati sensibili per assicurare la conformità, prevenire la perdita di dati e mantenere il controllo sulla proprietà intellettuale.

Con Netskope One DSPM, le aziende possono proteggere attivamente i dati sensibili, garantendo che l'addestramento dei modelli AI rimanga sicuro, conforme e controllato.



Sfida n. 4: Minacce mirate ai sistemi AI

Gli avversari stanno sviluppando tattiche per sfruttare le vulnerabilità specifiche dell'AI, usando prompt injection, data poisoning e input avversari progettati per distorcere i risultati o estrarre dati sensibili. Inoltre le applicazioni AI sono spesso integrate con sistemi aziendali più ampi, il che le rende un potenziale punto di accesso per movimenti laterali, escalation dei privilegi o furto di dati.

Non importa se un attore malevolo vuole manipolare l'output di un modello AI, estrarre dati di addestramento o sfruttare controlli di accesso deboli attorno alle API AI, la superficie di attacco si sta espandendo rapidamente. A complicare ulteriormente la faccenda è la mancanza di standard di sicurezza per proteggere i sistemi AI, il che lascia molte aziende impreparate a difendersi da nuovi vettori di attacco. Con il diffondersi dell'AI, cresce anche la necessità per i team di sicurezza di rilevare e mitigare attivamente le minacce mirate agli ambienti AI, prima che compromettano dati sensibili, operazioni o processi decisionali.



Come risolve il problema Netskope

Netskope affronta le crescenti minacce mirate ai sistemi AI con un approccio di sicurezza multilivello che integra protezione avanzata, visibilità approfondita e difese specifiche per l'AI.

Tra le capacità principali:

- **Protezione avanzata dalle minacce:** Usare l'apprendimento automatico, il sandboxing e l'analisi euristica per rilevare e bloccare sia le minacce conosciute che quelle zero-day, inclusi i malware nascosti nei file inviati agli strumenti AI.
- **Red Teaming e valutazioni delle vulnerabilità:** Netskope sta sviluppando attivamente capacità per testare i sistemi AI contro attacchi avversari, manomissioni dei modelli e perdite di dati per individuare le vulnerabilità prima che possano essere sfruttate.
- **Monitoraggio proattivo delle attività AI:** Rilevare minacce emergenti e vulnerabilità con il monitoraggio in tempo reale delle interazioni AI per assicurare una strategia di difesa completa.

Combinando queste tecnologie, Netskope offre una soluzione integrata che aiuta le aziende a proteggere i loro sistemi AI da minacce informatiche sofisticate e vettori di attacco in evoluzione.



Sfida n. 5: Esposizione dei dati

Una delle sfide più urgenti e pericolose nella sicurezza dell'AI è il rischio di esposizione dei dati. Man mano che i diversi dipartimenti adottano strumenti AI per aumentare la produttività, potrebbero caricare o condividere inconsapevolmente dati sensibili come codice sorgente, registri dei clienti, documenti finanziari o proprietà intellettuale riservata con modelli AI pubblici. Una volta esposti, questi dati possono essere conservati, usati per l'addestramento dei modelli o addirittura divulgati, in base alle politiche sulla privacy e alle pratiche di gestione dei dati dell'applicazione.

A differenza dei tradizionali canali di condivisione dei dati, le piattaforme di intelligenza artificiale possono agire da scatole nere, offrendo poca trasparenza sul modo in cui i dati vengono archiviati, consultati o usati. Senza barriere di protezione, le aziende affrontano rischi seri che vanno dalle violazioni normative al furto di proprietà intellettuale, ai danni d'immagine e allo svantaggio competitivo.

+ **I Netskope Threat Labs hanno osservato un'esposizione del codice sorgente in quasi il 50% delle violazioni delle politiche relative all'AI. Questo sottolinea quanto facilmente gli asset aziendali critici possono essere compromessi attraverso azioni apparentemente innocue, come incollare un frammento di codice in un chatbot AI per eseguire il debug o ottimizzarlo.**



Come risolve il problema Netskope

Netskope offre una protezione completa e ricca di contesto per i dati aziendali, sia a riposo che in movimento. Combinando valutazioni del rischio in tempo reale, controlli inline e basati su API, e verifiche della postura, le politiche di sicurezza unificate di Netskope consentono una governance precisa delle interazioni tra utenti e dati nell'azienda.

Tra le capacità principali:

- **Prevenzione avanzata della perdita di dati (DLP):** Proteggere le informazioni sensibili da strumenti AI come ChatGPT, Microsoft Copilot e Google Gemini, che possono estrarle ovunque si trovino gli utenti, in ufficio, a casa o altrove.
- **Controllo granulare:** Bloccare o limitare azioni ad alto rischio, come caricare il codice sorgente o documenti riservati, per prevenire il movimento non autorizzato dei dati.
- **Coaching degli utenti in tempo reale:** Formare gli utenti sulle violazioni delle politiche con suggerimenti visivi, per aiutare a ridurre le infrazioni ripetute e alleggerire il carico sui team di sicurezza.

Con queste capacità, Netskope assicura una protezione dei dati completa e adattiva che si estende all'intero ambiente AI e cloud di un'azienda.



Sfida n. 6: Governance, conformità e uso etico

Con il sempre più rapido diffondersi dell'AI, le aziende affrontano una crescente pressione per allinearsi agli standard di governance emergenti, ai requisiti normativi e alle aspettative etiche, soprattutto in settori altamente regolamentati come la finanza, la sanità e il governo. I Paesi di tutto il mondo stanno velocemente introducendo quadri e mandati specifici per l'AI, come si vede nel Regolamento sull'AI dell'UE, nel Quadro di gestione del rischio dell'AI del NIST e negli ordini esecutivi degli Stati Uniti sulla sicurezza dell'AI. Queste normative mirano ad assicurare uno sviluppo e un'implementazione responsabili dei sistemi AI, imponendo trasparenza, privacy dei dati, spiegabilità e non discriminazione.

Tuttavia, soddisfare questi standard non è un compito facile. I team di sicurezza e conformità devono comprendere come l'AI viene usata nel loro ambiente, garantire che i dati sensibili non vengano trattenuti o appresi in modo inadeguato e dimostrare l'adesione a linee guida legali ed etiche in evoluzione.



Come risolve il problema Netskope

Netskope assicura la governance dell'AI e la preparazione alla conformità attraverso la visibilità profonda, il controllo delle politiche e conoscenze in tempo reale sull'uso dell'AI in tutta l'azienda.

Tra le capacità principali:

- **Applicazione granulare delle politiche:** Controllare come vengono condivisi i dati con gli strumenti di AI, così da assicurare che i dati sensibili o regolamentati non vengano usati per l'addestramento non autorizzato di modelli di terze parti.
- **Controlli di conformità in tempo reale:** Bloccare il caricamento di informazioni sanitarie protette (PHI, *Protected Health Information*) su app non conformi o interrompere l'elaborazione dei dati finanziari in strumenti privi delle certificazioni adeguate.
- **Supporto al quadro normativo:** Facilitare la conformità con quadri come il Regolamento sull'AI dell'UE, il NIST AI RMF e altre normative in evoluzione.

Combinando visibilità, comprensione della conformità e applicazione adattiva delle politiche, Netskope consente alle aziende di abbracciare l'innovazione AI in modo responsabile, soddisfacendo le esigenze etiche e normative di oggi e di domani.

Il futuro della sicurezza dell'AI

Minacce e tecnologie emergenti

Con il rapido diffondersi dell'AI e l'emergere di nuovi casi d'uso, dai copiloti agli agenti AI personalizzati, il panorama delle minacce si sta evolvendo altrettanto rapidamente. Mentre gran parte dell'attenzione oggi si concentra sulla protezione dei dati e sull'integrità dei modelli, ci sono due aree emergenti dello sviluppo tecnologico pronte a presentare sfide ancora maggiori nel prossimo futuro.

In primo luogo, i sistemi di AI agentica, capaci di prendere decisioni e compiere azioni con un minimo di supervisione umana, sono in aumento. Secondo Gartner, entro il 2028, almeno il 15% delle decisioni aziendali quotidiane sarà preso autonomamente da AI agentica, rispetto a quasi lo 0% di oggi.⁴ Questo cambiamento aumenta in misura esponenziale la superficie di attacco, soprattutto se gli agenti ricevono l'accesso ai sistemi e ai dati aziendali.

In secondo luogo, l'AI fisica, come si vede nei veicoli e nei robot autonomi, sta guadagnando terreno in settori come la logistica, i trasporti e la produzione. Questi sistemi introducono rischi per la sicurezza nel mondo reale, dove un'AI compromessa o malfunzionante non causa solo perdita di dati, ma potenziale danno a persone e infrastrutture.

Man mano che le capacità dell'AI diventano più avanzate e profondamente integrate nelle operazioni aziendali, i leader della sicurezza devono stabilire una governance strategica e lungimirante. Alcune considerazioni chiave per rimanere un passo avanti:

- **Visibilità dell'uso dell'AI:** Conoscere quali team costruiscono o usano modelli AI, sia open che shadow IT. Garantire visibilità e supervisione centrale senza soffocare l'innovazione.
- **Affidabilità dei dati:** Garantire che i modelli siano addestrati su set di dati sicuri, conformi e ad alta integrità. Dati scadenti o contaminati portano a risultati inaccurati, distorti o non sicuri.
- **Autonomia e limiti di rischio:** Man mano che l'AI agentica diventa più capace, è bene definire chiare linee guida per l'autonomia. Non aspettare che gli agenti inizino a prendere decisioni ad alto impatto prima che la governance sia in atto.
- **Gestione del ciclo di vita del modello:** Trattare i modelli AI come codice: con controllo delle versioni, scansione delle vulnerabilità, controlli degli accessi e registri di audit.
- **Prontezza culturale:** La sicurezza non è solo tecnica, è comportamentale. È bene formare dipendenti e dirigenti sui rischi dell'AI, sull'uso sicuro e sul panorama normativo in evoluzione.

Il futuro della sicurezza dell'AI sarà definito non solo da quanto bene le aziende si proteggono dalle minacce odierne, ma da quanto attentamente si preparano per ciò che sta per arrivare.

⁴ Gartner 2024 <https://www.gartner.com/en/newsroom/press-releases/2024-10-21-gartner-identifies-the-top-10-strategic-technology-trends-for-2025>

Predizione

La società di analisi Gartner prevede che, entro il 2028, almeno il 15% delle decisioni aziendali giornaliere verrà preso autonomamente attraverso l'AI agentica, rispetto allo 0% nel 2024.

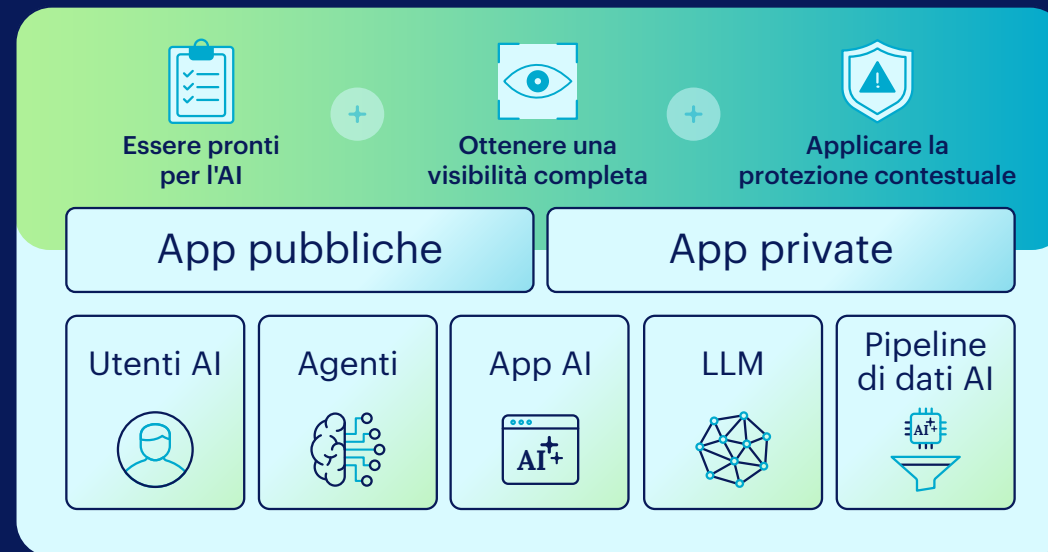
Conclusioni

Sicurezza dell'AI end-to-end, ovunque con Netskope One

Mentre le imprese si affrettano ad adottare l'AI, i leader della sicurezza affrontano una pressione sempre più alta per proteggere i dati sensibili e rimanere un passo avanti rispetto ai nuovi rischi che mirano al loro ecosistema AI. Dalla mancanza di visibilità nell'uso dell'AI all'esposizione dei dati e alle esigenze di conformità, abbiamo delineato sei sfide fondamentali che i team di sicurezza devono affrontare per abilitare in modo sicuro l'AI in tutta l'impresa:

- Mancanza di visibilità
- Capire il rischio delle applicazioni AI
- Integrità del modello AI
- Minacce mirate ai sistemi AI
- Esposizione dei dati
- Governance, conformità e uso etico

Netskope può aiutare a risolvere questi problemi con Netskope One Security Service Edge (SSE). Netskope One SSE protegge da minacce avanzate e abilitate al cloud e salvaguarda i dati su tutti i vettori (cloud, app, utente). I nostri strumenti offrono una riduzione misurabile del rischio, oltre a un minor volume di incidenti e tempi medi di risoluzione più rapidi.



Ricerca

La società di analisi Forrester ha scoperto che Netskope offre una riduzione dell'80% del rischio di una violazione grave causata da un attacco esterno, pari a un risparmio di 2 milioni di dollari in costi annualizzati per violazioni gravi.⁵

⁵ Rapporto Forrester: The Total Economic Impact™ of Netskope SSE

<https://www.netskope.com/resources/analyst-reports/forrester-the-total-economic-impact-of-netkskope-sse>

A proposito di Netskope

Netskope, leader nella sicurezza e nelle reti moderne, soddisfa le esigenze dei team di sicurezza e networking fornendo accesso ottimizzato e sicurezza in tempo reale basata sul contesto per persone, dispositivi e dati ovunque si spostino. Migliaia di clienti, fra cui più di 30 appartenenti a Fortune 100, si affidano alla piattaforma Netskope One, al suo Zero Trust Engine e alla sua potente rete NewEdge per ridurre i rischi e ottenere massima visibilità e controllo su applicazioni cloud, AI, SaaS, Web e private che offrono sicurezza e prestazioni accelerate senza compromessi.

Vuoi saperne di più?

[Chiedi una demo](#)



©2025 Netskope, Inc. Tutti i diritti riservati. Netskope, NewEdge, SkopeAI e il logo "N" sono marchi registrati di Netskope, Inc. Netskope Active, Netskope Cloud XD, Netskope Discovery, Cloud Confidence Index e SkopeSights sono marchi di Netskope, Inc. Tutti gli altri marchi inclusi sono marchi dei rispettivi proprietari. 06/25 EB-827-1-IT

Risorse



**Protezione dell'AI con
Netskope One**



**Dominare l'adozione dell'AI
con sicurezza end-to-end,
ovunque - Blog**



**Netskope Threat Labs:
Rapporto delle minacce
al cloud dell'AI generativa**



**Protezione dell'AI
generativa per principianti**



I

01

02

03

04

C