

GUIDA PER PMI · L'EVOLUZIONE DI WAZUH

Agentic SOC per PMI

Come un analista AI ti fa vivere gli alert giusti, mentre l'umano resta al comando.

Una guida operativa per chi gestisce Wazuh in una PMI: come passare dal rumore di migliaia di alert a pochi incidenti che una persona può davvero seguire, senza mai rinunciare al controllo umano sulle decisioni.

IN QUESTA GUIDA

- 01  La crisi degli alert
- 02  Agentic AI senza fuffa
- 03  L1 AI su ogni alert
- 04  Da 1000 alert a 3 incidenti
- 05  Human-in-the-loop



wazuh

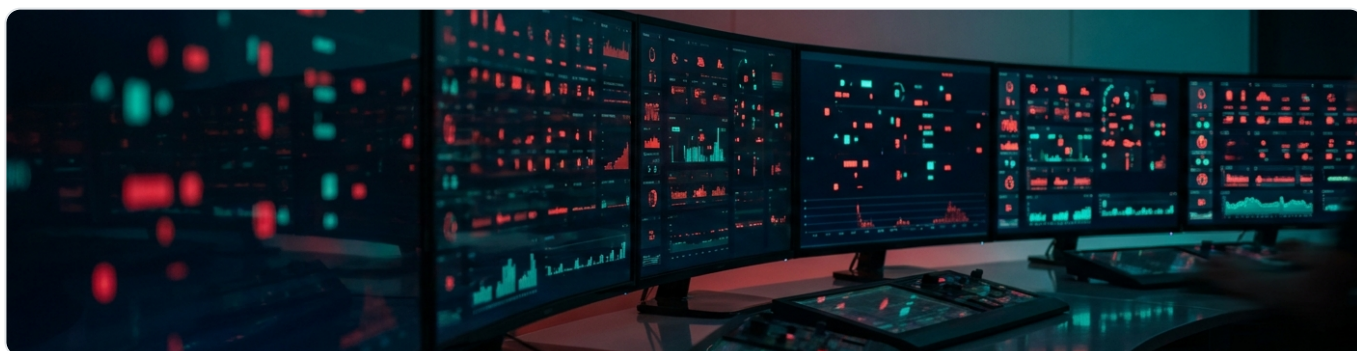
ArcaSafe S.r.l. · Partner Wazuh · Azienda certificata ISO/IEC 27001 e ISO 9001
SOC gestito on-prem con tecnici certificati · Cybersecurity, NIS2 & GDPR

✱ Ogni giorno il tuo Wazuh genera più segnali di quanti una persona possa leggerne. Non è un problema di strumenti: è un problema di scala. Questa guida spiega, senza fuffa, come un SOC agentico trasforma migliaia di alert in pochi incidenti gestibili, tenendo la decisione finale dove deve stare: nelle mani di chi risponde.

CAPITOLO 01

⚠ La crisi degli alert

Il paradosso della sicurezza moderna è che più strumenti installi, più rumore produci. Un deployment Wazuh mediamente configurato in una PMI può emettere migliaia di alert al giorno: login falliti, processi anomali, modifiche a file, connessioni verso IP sconosciuti. Ognuno, preso singolarmente, potrebbe essere qualcosa. Nessuno, da solo, ti dice cosa sta davvero succedendo.



Il volume di alert è oggi il vero vettore di rischio.

Questo genera la cosiddetta alert fatigue, e non è un fastidio: è un rischio operativo concreto. Quando un analista rivede il duecentesimo falso positivo della giornata, la sua soglia di attenzione crolla. È statistica, non pigrizia. Ed è esattamente lì che l'attacco vero passa inosservato, mimetizzato tra decine di segnali innocui che nessuno ha avuto tempo di leggere.

Per una PMI il problema si aggrava per un motivo strutturale: il team è piccolo. Spesso è una persona, a volte mezza persona, che fa sicurezza tra un ticket IT e l'altro. La reazione tipica — alzare le soglie, silenziare regole — abbassa il rumore ma anche la copertura: una toppa che sposta il rischio, non lo elimina.

La via d'uscita non è produrre meno segnali. È metterci in mezzo qualcosa che li legga tutti, sempre, senza stancarsi, e che porti in superficie solo ciò che merita l'attenzione di un essere umano. Questo qualcosa è un analista AI.

✓ **Il volume è il vero rischio**

Non un dettaglio da gestire "quando avremo tempo": è il vettore di rischio principale.

CAPITOLO 02

🔗 Agentic AI senza fuffa

Agentic AI è diventato un termine da slide commerciale, quindi mettiamo dei paletti concreti. Un agente non è un chatbot che risponde a domande, e non è un modello che indovina una risposta. Un agente è un sistema che percepisce una situazione, ragiona su di essa, usa strumenti per verificare, e compie un'azione. Quattro componenti, tutti necessari.



Reasoning, memoria, strumenti e azione: i quattro componenti di un vero agente.

Il reasoning è la capacità di seguire un ragionamento strutturato. La memoria attinge a un corpus di playbook — procedure e pattern noti — applicandoli al caso concreto: è la differenza tra un neoassunto e un analista esperto. Gli strumenti gli permettono di verificare davvero: un agente che genera decoder dai log grezzi, li testa con logtest e si autocorregge in minuti invece di ore.

E qui il punto che distingue un sistema serio da una demo. La filosofia giusta non è l'autonomia, è il coordinamento. Un agente autonomo che decide tutto da solo sbaglia in silenzio e su scala. Un agente ben progettato conosce i propri limiti: quando è sicuro procede, quando è incerto si ferma ed escala a un umano.

✓ I 4 componenti di un vero agente

Reasoning · Memoria (playbook) · Strumenti · Azione. Tutti necessari.

CAPITOLO 03

🔗 L1 AI su ogni alert

Il primo livello di un SOC, l'L1, guarda ogni singolo alert e decide se merita attenzione. In un team umano è il collo di bottiglia: nessuno riesce a guardarli tutti, quindi si campiona e si spera che il campione contenga le cose giuste.



Verdetto, risk score, MITRE ATT&CK e motivazione scritta, per ogni singolo alert.

Un analista AI L1 rimuove il compromesso: guarda ogni alert Wazuh, senza eccezioni e senza stancarsi. Ma la parte interessante è la qualità del prodotto che restituisce. Per ciascun alert: un verdetto (benigno, sospetto, malevolo); un risk score che ordina la coda per rischio reale e non per la severità nominale della regola; un arricchimento MITRE ATT&CK che traduce l'evento nel linguaggio delle tecniche; e una motivazione scritta.

Quest'ultimo punto cambia la natura del rapporto uomo-macchina. Il tutto attingendo a un corpus di play-book: conoscenza consolidata applicata caso per caso, non improvvisazione. La coda mattutina non è più un muro di segnali, ma una lista ordinata di verdetti motivati.

✓ Per ogni alert, quattro cose

Verdetto · Risk score · Mappatura MITRE ATT&CK · Motivazione scritta.

CAPITOLO 04

↔ Da 1000 alert a 3 incidenti

Anche il miglior L1 lascia un problema: la frammentazione. Un attacco reale non produce un alert, ne produce decine — phishing, esecuzione, persistenza, movimento laterale — su host e momenti diversi. Vederli come mille eventi separati significa non vederli affatto.



La correlazione L2 fonde lo sciame di duplicati in poche storie coerenti.

Qui entra il secondo passaggio, l'L2, ed è ciò che separa uno strumento di triage da un vero SOC. L'L2 non guarda l'alert singolo: guarda le relazioni. Correla le escalation dell'L1 e le fonde in un incidente — una storia coerente con inizio, sviluppo e portata — invece di uno sciame di duplicati.

È da qui che nasce il 1000 a 3, e non è marketing: è la conseguenza matematica del correlare invece di elencare. Un carico che una persona sola, in una PMI, può davvero gestire — e che rende possibile rispettare i tempi di notifica NIS2.

✓ Il nodo NIS2

Primo alert entro 24 ore: solo un incidente già strutturato rende la notifica tempestiva possibile.

CAPITOLO 05

Human-in-the-loop

C'è la tentazione, quando l'automazione funziona, di togliere l'umano di mezzo. È l'errore più costoso. Un SOC agentico serio non punta alla piena autonomia: mette la persona giusta davanti alla decisione giusta, con il lavoro preparatorio già fatto.



L'AI toglie la fatica; l'umano mantiene giudizio, contesto e decisione.

Alcune risposte partono in automatico — il blocco di un IP palesemente malevolo via Wazuh Active Response, dove ogni minuto conta. Altre passano per un'approvazione esplicita. Perché tenere l'umano al comando

è corretto? Contesto di business (l'AI non sa che quel server sta chiudendo la contabilità), responsabilità tracciabile, e gestione dell'incertezza: quando l'agente è dubbioso, escala invece di tirare a indovinare.

Il risultato non è AI contro umano, ma una divisione del lavoro onesta — il modello che permette a una PMI la copertura di un SOC strutturato senza raddoppiare il team, e che regge in ambito NIS2 perché la direttiva chiede governance e responsabilità, non deleghe cieche.



Chi decide, decide

Azioni automatiche dove il rischio è basso; approvazione esplicita dove serve giudizio.

Vuoi vederlo sui tuoi numeri?

Il punto non è comprare l'ennesimo strumento, ma cambiare il rapporto tra il tuo team e il rumore che lo sommerge. In ArcaSafe implementiamo il SOC agentic come servizio gestito on-prem, in alta affidabilità, con un tecnico certificato al tuo fianco e l'integrazione con l'incident response.

Richiedi una call sul SOC Wazuh →

Un click apre un'email già compilata (oggetto Wazuh): aggiungi i tuoi dati e invia. Oppure scrivici a commerciale@arcasafe.eu.

ArcaSafe S.r.l. · Via G. Segantini 5, 20825 Barlassina (MB) · CF e P.IVA 09546870966 · info@arcasafe.eu · PEC arcasafe@pec.it · arcasafe.eu
· Partner Wazuh · ISO/IEC 27001 e ISO 9001