

AI - SEC

Secure GenAI e Agent Threat Modeling

Due giorni, dopo i quali Lei ha un threat model, un diagramma del flusso di dati e un piano di test per un'applicazione reale invece di una lista di rischi teorici.

DURATA

2 giorni

GRUPPO

fino a 12 persone

FORMATO

In-house · Remoto · Nel lab · Tedesco o inglese

QUESTO ESISTE DOPO

- ▶ Un diagramma del flusso di dati della Sua applicazione, dall'input attraverso modello, fonti di sapere e strumenti fino all'output
- ▶ Un threat model che trasferisce una sistematica in stile STRIDE su GenAI e agenti
- ▶ Un catalogo di misure di protezione concrete: Guardrails, permessi e filtri di output, assegnati per ogni minaccia
- ▶ Un piano di test che rende verificabile se una misura funziona
- ▶ Una lista dei rischi residui con priorità e responsabili

AGENDA

Giorno 1 · 09:00	Perché la GenAI ha minacce diverse: Prompt Injection, output non sicuri, fuga di dati attraverso fonti di sapere e strumenti. Domanda centrale: dove questo si distingue dalla sicurezza applicativa classica e dove no?
Giorno 1 · 10:30	Diagramma del flusso di dati della Sua applicazione: Disegniamo il flusso completo: input, modello, fonti RAG, chiamate a tool, output. Esercizio: ogni trust boundary viene marcato e nominato.
Giorno 1 · 13:30	Trasferire una sistematica in stile STRIDE sulla GenAI: Percorriamo le categorie di minaccia e le traduciamo su prompt, contesto, strumenti e output. Il quadro di riferimento OWASP LLM serve da checklist, perché nessuna categoria venga dimenticata.
Giorno 1 · 15:30	Raccogliere le minacce sulla Sua applicazione: Esercizio sul Suo diagramma del flusso di dati: per ogni trust boundary nominiamo attacchi plausibili. Risultato del giorno: un primo threat model con minacce in ordine di priorità.
Giorno 2 · 09:00	Dimostrazione dal vivo nel Lab: Mostriamo in un ambiente preparato come agisce una Prompt Injection attraverso una fonte RAG e cosa rivela uno strumento con permessi troppo ampi. Nessun dato di clienti, solo per capire l'effetto.
Giorno 2 · 11:00	Misure di protezione per ogni minaccia: Guardrails, permessi secondo il minimo privilegio, filtri di output, verifica delle fonti. Decisione per ogni minaccia: quale misura, quale sforzo, quale rischio residuo resta?
Giorno 2 · 14:00	Costruire il piano di test: Per ogni misura un test che ne dimostra l'effetto. Esercizio: per ogni minaccia importante formuliamo un caso di test concreto con risultato atteso.
Giorno 2 · 16:00	Rischi residui e consegna: Raccogliamo ciò che resta e nominiamo i responsabili. Risultato del giorno: threat model, catalogo delle misure, piano di test e lista dei rischi residui con priorità per la Sua applicazione.



IL SUO TRAINER

Dino Bordonaro

Microsoft MVP per otto anni consecutivi (Microsoft Azure · Cloud and Datacenter Management), Microsoft Certified Trainer, quasi 30 anni di IT. Costruisce sistemi AI sovrani su hardware proprio (Azure Local, Foundry Local, GPU NVIDIA) e forma esattamente su questo. Oltre 800 professionisti formati, ogni concetto collaudato nel proprio enterprise lab in un datacenter certificato TÜV.

15.800 € forfettario, IVA esclusa

Il track fornisce modello, misure e piano di test come bozza. Non sostituisce un'implementazione, né un mandato di penetration test, né un servizio di sicurezza gestito.

AI - SEC