

¿Por qué DLP, los firewalls de IA y los LLM privados no logran detener las fugas de datos con IA?

Respuesta breve: Cada uno aborda parte del problema desde la capa equivocada. El DLP puede eludirse, los firewalls de IA se sitúan en el borde de la red y son ciegos a la actividad en el dispositivo, y los LLM privados son costosos sin resolver nada respecto a los empleados que igualmente usan herramientas públicas.

Dónde falla cada enfoque

- Los **LLM privados** son costosos de crear y mantener, y no impiden que nadie abra ChatGPT en una pestaña del navegador de todos modos.
- Los **firewalls de IA** inspeccionan el tráfico de red, lo que los deja ciegos ante lo que ocurre en el propio endpoint y ante la actividad interna.
- El **DLP tradicional** fue diseñado para archivos y flujos de correo electrónico, y es eludido rutinariamente por vías que nunca fue concebido para vigilar.

El problema de la lógica invertida

Varios enfoques de esta categoría intentan proteger los datos frente a la IA insertando *otra* IA para filtrar la salida. Esto añade costo y una nueva dependencia sin abordar el punto real donde los datos salen: el endpoint, en el momento del acceso.

FortressAI aplica su protección a nivel del kernel, por debajo de todas estas capas, y se aplica sin importar qué herramienta, navegador o extensión esté involucrada.

Revision #1

Created 2026-07-24 01:50:07 UTC by Dennis Underwood

Updated 2026-07-24 01:50:07 UTC by Dennis Underwood