

Pourquoi la DLP, les pare-feux IA et les LLM privés ne parviennent-ils pas à empêcher les fuites de données liées à l'IA ?

Réponse courte : Chacune de ces solutions traite une partie du problème depuis la mauvaise couche. La DLP peut être contournée, les pare-feux IA se situent en périphérie du réseau et sont aveugles à l'activité sur l'appareil, et les LLM privés coûtent cher tout en ne faisant rien contre les employés qui utilisent malgré tout des outils publics.

Où chaque approche est insuffisante

- Les **LLM privés** sont coûteux à construire et à exploiter, et n'empêchent en rien quiconque d'ouvrir ChatGPT dans un onglet de navigateur.
- Les **pare-feux IA** inspectent le trafic réseau, ce qui les rend aveugles à ce qui se passe sur le point de terminaison lui-même et à l'activité interne (insider).
- La **DLP traditionnelle** a été conçue pour les fichiers et les flux de messagerie, et est régulièrement contournée par des voies qu'elle n'a jamais été conçue pour surveiller.

Le problème de la logique inversée

Plusieurs approches de cette catégorie tentent de protéger les données contre l'IA en insérant une *autre* IA pour filtrer la sortie. Cela ajoute un coût et une nouvelle dépendance sans traiter l'endroit où les données sortent réellement — le point de terminaison, au moment de l'accès.

FortressAI applique ses règles au niveau du noyau, en dessous de toutes ces couches, et s'applique quel que soit l'outil, le navigateur ou l'extension concerné.

Revision #1

Created 2026-07-24 06:59:07 UTC by Dennis Underwood

Updated 2026-07-24 06:59:07 UTC by Dennis Underwood