

Os hackers usarão mais IA em ataques de ransomware em 2025?

Resposta curta: Sim. Espera-se que os invasores continuem a expandir o uso de inteligência artificial, especialmente modelos de linguagem de grande escala, para se passar por indivíduos confiáveis e manipular vítimas a conceder acesso ou confiança.

Por Que a IA Se Tornou a Favorita dos Hackers

Resultados de pesquisa, feeds de notícias e conversas cotidianas estão agora saturados de referências à IA, e esse mesmo entusiasmo se instalou entre os cibercriminosos. Quando a maioria das pessoas menciona IA, normalmente está se referindo a modelos de linguagem de grande escala, como aqueles por trás do ChatGPT e de ferramentas semelhantes, capazes de imitar a comunicação humana de forma convincente. Os invasores reconheceram que essa capacidade é extremamente útil para a engenharia social. Em vez de depender de e-mails de phishing genéricos, eles podem gerar mensagens, vozes ou conversas que se assemelham muito a um colega, fornecedor ou figura de autoridade real, tornando muito mais fácil conquistar a confiança de um alvo antes de atacar.

Esprete Táticas de Personificação Mais Refinadas

A tendência para 2025 não é que os ataques baseados em IA sejam algo totalmente novo, mas sim que estão se tornando mais sofisticados e profissionais. Da mesma forma que empresas legítimas estão implementando agentes de vendas com IA e ferramentas automatizadas de prospecção, os invasores estão refinando suas próprias técnicas de personificação baseadas em IA. Isso significa mensagens mais convincentes que imitam chefes, colegas de trabalho ou parceiros confiáveis, todas projetadas para explorar as relações de trabalho nas quais as pessoas confiam diariamente. O perigo é que uma conversa que pareça pessoal e legítima pode, na verdade, estar vindo de um sistema de IA altamente capaz, construído para manipular, não para ajudar.

O Que Isso Significa para a Defesa

À medida que essas táticas de personificação amadurecem, as organizações devem presumir que os ataques baseados em confiança se tornarão mais difíceis de identificar apenas por instinto. Verificar solicitações por meio de canais independentes, em vez de confiar em uma mensagem ou chamada à primeira vista, torna-se cada vez mais importante. Como esses ataques são projetados para contornar o julgamento humano, contar com capacidades automatizadas de detecção e resposta, como as incorporadas ao FortressAI da Cyber Crucible, oferece uma camada adicional de proteção quando a engenharia social consegue ultrapassar as defesas de uma pessoa.

<https://player.vimeo.com/video/1046827321>

[Assistir no Vimeo](#) · Legendas: English, Français, Español, العربية

Revision #1

Created 2026-07-24 06:15:48 UTC by Dennis Underwood

Updated 2026-07-24 06:15:49 UTC by Dennis Underwood