
Um quarto das empresas do Reino Unido e dos EUA sofrem ataques de en

Data: 2025-09-17 11:15:00

Autor: Inteligência Against Invaders

Os líderes britânicos e americanos de segurança cibernética estão cada vez mais preocupados com a expansão da superfície de ataque de IA, particularmente o uso não autorizado de ferramentas de IA e as tentativas de corromper dados de treinamento, de acordo com uma nova pesquisa da IO.

O especialista em segurança e conformidade entrevistou 3000 líderes de segurança de TI em ambos os lados do Atlântico para compilar seu terceiro relatório anual *Relatório do Estado da Segurança da Informação*, que foi publicado esta manhã.

Ele revelou que pouco mais de um quarto (26%) sofreu um ataque de envenenamento de dados, que ocorre quando os agentes de ameaças procuram interferir nos dados de treinamento do modelo para alterar seu comportamento.

Esses ataques podem ser lançados para sabotar organizações que dependem de modelos de IA ou então apoiar os agentes de ameaças de maneiras mais direcionadas, como fazer com que os sistemas de detecção de malware falhem.

Até agora, os ataques de envenenamento de dados eram considerados mais teóricos do que generalizados.

[Leia mais sobre ameaças de IA: Escassez de talentos afeta 80% das empresas do Reino Unido atingidas por ameaças de IA](#)

O relatório da IO também revelou que 37% das empresas estão vendo os funcionários usarem ferramentas de IA generativa (GenAI) na empresa sem permissão.

Esse tipo de shadow AI pode apresentar grandes riscos associados ao vazamento de dados e violações de conformidade, bem como possíveis vulnerabilidades se a ferramenta GenAI em questão não for segura.

O carro-chefe da DeepSeek, o LLM R1, foi encontrado no início deste ano [para conter várias vulnerabilidades](#). A empresa também [exposto acidentalmente](#) um banco de dados de históricos de bate-papo e outras informações confidenciais do usuário.

Preocupações e confiança no futuro

Os entrevistados do relatório pareciam em conflito com suas atitudes em relação à IA. Por um lado,

eles citaram as maiores ameaças emergentes à segurança cibernética para o próximo ano como phishing gerado por IA (38%) e desinformação (42%), shadow AI (34%) e falsificação de identidade em reuniões virtuais (28%).

No entanto, os incidentes de ataques relacionados a deepfake caíram de 33% no ano passado para 20%, de acordo com a IO.

Além disso, os entrevistados pareciam otimistas em relação ao futuro. A grande maioria disse que se sente “preparada” para se defender contra phishing gerado por IA (89%), falsificação de identidade (84%), malware orientado por IA (87%) e desinformação (88%), shadow AI (86%) e envenenamento de dados (86%).

Três quartos (75%) [dito](#) eles estão implementando políticas de uso aceitáveis para IA, o que deve pelo menos ajudar a mitigar o uso não autorizado de ferramentas.

Chris Newton-Smith, CEO da IO, descreveu a IA como uma faca de dois gumes.

“Embora ofereça uma enorme promessa, os riscos estão evoluindo tão rápido quanto a própria tecnologia. Muitas organizações correram e agora estão pagando o preço”, acrescentou.

“Os ataques de envenenamento de dados, por exemplo, não apenas prejudicam os sistemas técnicos, mas ameaçam a integridade dos serviços dos quais dependemos. Adicione a IA sombra à mistura e fica claro que precisamos de uma governança mais forte para proteger as empresas e o público.”