

Relatório: Do risco dos agentes às conquistas humanas

Introdução

O cenário corporativo ultrapassou a fase de experimentação teórica com IA e organizações brasileiras já começaram a incorporar agentes autônomos e automações baseadas em IA em workflows corporativos.

No entanto, enquanto as organizações correm para capturar ganhos de eficiência, suas posturas de segurança estão se fragmentando. O relatório ["From Agentic Risk to Human Wins"](#), da KnowBe4, destaca os pontos cegos críticos, vulnerabilidades psicológicas e lacunas de preparação com as quais as organizações do país estão lidando.

Seção 1: Shadow AI e o Déficit de Governança

Enquanto organizações introduzem IA em suas stacks formais de tecnologia, uma camada significativa e não gerenciada de "Shadow AI" (*uso de ferramentas como chatbots, geradores de imagem ou automações sem o conhecimento, autorização ou supervisão da área de TI e de segurança da empresa*) está emergindo entre as forças de trabalho. Agentes maliciosos sabem muito bem que equipes de segurança não conseguem proteger ativos que não conseguem enxergar.

- **65%** das organizações pesquisadas utilizam agentes autônomos de IA capazes de agir por conta própria, porém **60%** dos líderes admitem que o uso de IA dentro de seus ambientes é totalmente não aprovado ou carece de governança corporativa formal.
- **63%** dos tomadores de decisão em cibersegurança afirmam que o uso não autorizado de softwares externos e aplicações de IA não aprovadas esteve entre os comportamentos que mais impactaram a cibersegurança de suas organizações nos últimos 12 meses.
- **47%** dos trabalhadores brasileiros que utilizam IA no trabalho afirmam recorrer, ao menos ocasionalmente, a ferramentas de IA escolhidas por conta própria, em vez de utilizar exclusivamente as soluções disponibilizadas pela organização.

Destaques

- ▶ **Descubra a lacuna de governança:** Entenda a dimensão do Shadow AI e por que o uso de IA nas organizações frequentemente ocorre sem uma governança corporativa formal.
- ▶ **Avalie a mudança para agentes autônomos:** Descubra o percentual de organizações que já utilizam agentes de IA autônomos capazes de agir por conta própria e as consequências de segurança associadas a essa adoção.
- ▶ **Conheça o campo de batalha psicológico:** Entenda como o vetor de ameaça evoluiu para deepfakes hiper-realistas criados por IA e por que a maioria dos funcionários admite que poderia ser enganada com sucesso.
- ▶ **Identifique a verdadeira ameaça humana:** Veja os dados que mostram como a sobrecarga cognitiva e os obstáculos operacionais estão levando funcionários a contornar processos, contribuindo para uma parcela significativa dos incidentes de segurança.
- ▶ **Avalie sua preparação:** Compreenda o "Padrão Ouro" de maturidade em segurança e descubra se os líderes de segurança se consideram resilientes o suficiente para enfrentar as ameaças emergentes impulsionadas por IA.

Seção 2: Deepfakes e o Campo de Batalha Psicológico

O vetor da ameaça mudou fundamentalmente do phishing tradicional para manipulações hiper-realistas criadas por IA. As defesas cognitivas humanas estão falhando diante de conteúdos gerados por máquinas.

- Impressionantes **90%** dos funcionários da região afirmam que conteúdos de voz e vídeo criados com deepfake estão tão realistas que ficou mais difícil saber no que confiar.
- **61%** da força de trabalho local reconhece abertamente que poderia ser enganada com sucesso por um golpe sofisticado de deepfake se passando por um executivo ou stakeholder interno da empresa.
- O cenário de ameaças migrou da velocidade humana para ataques automatizados em larga escala. As organizações ainda demonstram baixa maturidade operacional para ataques automatizados e multietapas que atingem funcionários simultaneamente por e-mail, SMS e aplicativos de colaboração.

Seção 3: Fricção Operacional e a Lacuna do “Padrão Ouro”

A principal ameaça não é apenas o ataque avançado, é o efeito acumulado da fadiga humana em ambientes corporativos de alta pressão colidindo com uma baixa maturidade operacional em segurança.

- **62%** dos funcionários brasileiros admitem que restrições de tempo, sobrecarga de trabalho e distrações do dia a dia contribuem diretamente para erros de segurança, mesmo quando conhecem os protocolos adequados. Além disso, **43%** dos líderes apontam distrações no ambiente de trabalho como uma das principais causas de falhas humanas relacionadas à segurança.
- Líderes de cibersegurança no Brasil confirmam a gravidade desse desafio comportamental. Os dados sugerem que falhas operacionais cotidianas continuam sendo uma das principais fontes de risco: **65%** dos líderes apontam erros

não intencionais durante atividades normais de trabalho entre os comportamentos que mais impactaram a cibersegurança de suas organizações

- Apesar da crescente conscientização sobre segurança digital, apenas **33%** das organizações alcançaram um modelo de maturidade considerado “padrão ouro” — definido como uma cultura totalmente sincronizada, capaz de gerenciar simultaneamente riscos humanos e riscos associados à IA.
- Além disso, **95%** dos líderes de segurança no Brasil afirmam sentir que suas organizações estão preparadas para enfrentar ameaças emergentes ou inesperadas impulsionadas por inteligência artificial ao longo dos próximos 12 meses.

Conclusão

A força de trabalho no Brasil mudou. Funcionários e agentes de IA agora operam como uma única camada interconectada de defesa. Isso muda completamente o que significa uma boa segurança.

Organizações que alinham conscientização, comportamento e cultura em torno de uma visão clara de seus riscos fazem mais do que simplesmente gerenciar uma exposição. Elas se tornam alvos mais difíceis.

As ameaças não vão desaparecer — mas irão encontrar organizações mais preparadas.

