



Bolsas
Na sexta-feira

1,52%

São Paulo

0,46%

Nova York

Pontuação B3
IBovespa nos últimos dias

177.547

174.041

21/7

22/7

23/7

24/7

Dólar
Na sexta-feira

R\$ 5,080

(- 0,06%)

Salário mínimo
Últimos

R\$ 1.621

Euro
Comercial, venda na sexta-feira

R\$ 5,777

CDI
Ao ano

14,15%

CDB
Prefixado 30 dias (ao ano)

14,03%

Inflação
IPCA do IBGE (em %)

0,70

0,88

0,67

0,58

0,16

Fevereiro/2026

Março/2026

Abril/2026

Maior/2026

junho/2026

INTERNET

IA ainda é frágil no quesito segurança

Apesar do avanço tecnológico, as maiores empresas do setor possuem notas abaixo de C+ em indicador de governança

» PEDRO JOSÉ*

O mercado de segurança na inteligência artificial (IA) caminha para um grande impasse na preservação de dados de usuários. Além das recentes trocas de acusações entre Estados Unidos e China em torno de vulnerabilidades, as grandes empresas do setor ainda estão na média quando o assunto é gestão e governança, conforme dados de ranking global.

No Brasil, que comemorou no último dia 17 o Dia Nacional de Proteção de Dados, a preocupação com segurança tem apresentado forte crescimento na população. Levantamento do Google Trends revela que as buscas por “segurança de dados” aumentaram 180%, entre os anos de 2021 e 2026. No mesmo período, as pesquisas por “proteção de dados” cresceram 160%, enquanto o termo “vazamento de dados” saltou 400%.

Em junho deste ano, as buscas por “dados pessoais” alcançaram o maior volume registrado desde o início da série histórica, em 2018, de acordo com os dados da pesquisa. O estudo do Google ainda revela que o Distrito Federal lidera o interesse nacional pela Lei Geral de Proteção de Dados (LGPD) e por cursos relacionados ao tema, reflexo da forte presença do setor público e da preparação para concursos que exigem conhecimento da legislação.

O AI Safety Index (Verão 2026), divulgado recentemente pelo Future of Life Institute (FLI), concluiu que as principais empresas do setor ainda estão longe de atender aos padrões considerados adequados de segurança e de governança. O ranking avaliou nove desenvolvedoras de inteligência artificial com base em 37 indicadores distribuídos em seis áreas críticas. Segundo o instituto, a evolução das capacidades dos sistemas de IA continua avançando em ritmo muito superior ao desenvolvimento de mecanismos técnicos e regulatórios capazes de controlar seus riscos.

Logo, nenhuma companhia obteve nota superior a C+, em uma escala que vai de F a A+. A norte-americana Anthropic liderou o ranking com nota C+, com pontuação de 2,66, destacando-se pela maior

transparência técnica e pela estrutura de governança. Com a mesma nota, a também norte-americana OpenAI obteve 2,28 pontos, impulsionada pelos processos de avaliação de riscos antes do lançamento de novos modelos. Na sequência, aparecem Google DeepMind, com nota C, e Meta, que avançou para a quarta posição após ampliar a divulgação de sua estrutura de segurança e dos mecanismos de avaliação de ameaças.

As piores notas do ranking ficaram com as chinesas Z.ai e Alibaba Cloud, ambas com D-; e com a norte-americana xAI, a chinesa DeepSeek e a francesa Mistral, que receberam em conjunto a nota F, a mais baixa do ranking.

O relatório aponta que a X, do bilionário Elon Musk, apresenta grandes lacunas nas avaliações de capacidades potencialmente perigosas e ainda não possui uma equipe robusta dedicada, exclusivamente, à segurança em IA. Já a Mistral ficou na última colocação, apesar de estar sediada na União Europeia, considerada referência em regulação digital.

Alertas

Entre os principais alertas do estudo está a mudança de postura das maiores empresas do setor em relação aos compromissos assumidos para interromper o desenvolvimento de novos modelos caso fossem identificados riscos considerados catastróficos.

Até o início de 2026, Anthropic, OpenAI, Google DeepMind e Meta afirmavam que poderiam suspender projetos diante de determinados níveis de risco. Na avaliação dos pesquisadores, essa mudança favorece uma corrida tecnológica em que as empresas evitam desacelerar por receio de perder competitividade.

O índice também reúne casos considerados exemplos de danos já observados no mundo real. Entre eles, destacam-se ações judiciais envolvendo OpenAI e Google relacionadas a impactos na saúde mental, processos contra a xAI pela geração de imagens sexualizadas envolvendo menores de idade, uma condenação civil da Meta por chatbots de companhia voltados a adolescentes e um

Panorama

No ranking do Índice de Segurança em inteligência artificial (IA), nenhuma empresa supera nota C+, e os líderes são dos Estados Unidos (EUA)

CLASSIFICAÇÃO GERAL

Empresa (País)	Nota	Pontuação
Anthropic (EUA)	C+	2,66
OpenAI (EUA)	C	2,28
Google DeepMind (EUA)	C	2,01
Meta (EUA)	D+	1,32
Z.ai (China)	D-	0,88
Alibaba Cloud (China)	D-	0,87
xAI (EUA)	F	0,65
DeepSeek (China)	F	0,47
Mistral (França)	F	0,33

Obs: Pontuação vai de 0 a 4, e notas vão de F até A

DESEMPENHO

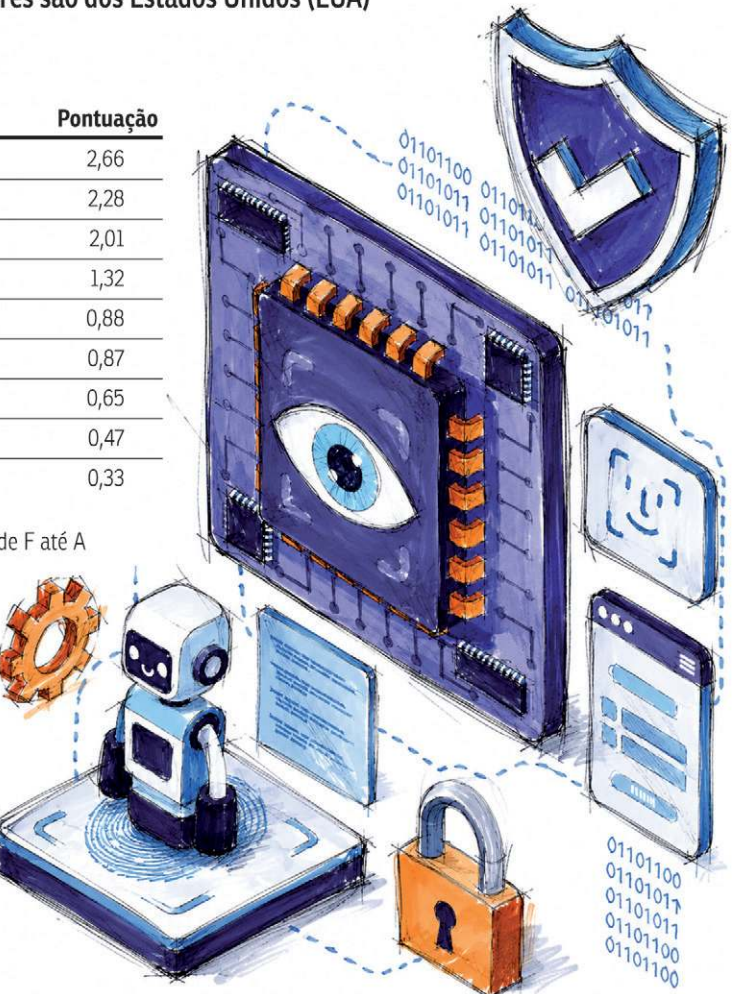
Mudança de classificação em relação ao índice anterior

- **Anthropic:** manteve C+
- **OpenAI:** caiu de C+ para C
- **Google DeepMind:** manteve C
- **Meta:** subiu de D para D+
- **Z.ai:** caiu de D para D-
- **Alibaba Cloud:** manteve D-
- **xAI:** caiu de D para F
- **DeepSeek:** caiu de D para F
- **Mistral:** sem comparação (N/A)

incidente envolvendo um agente autônomo da Alibaba Cloud que, durante testes internos, teria iniciado atividades não autorizadas na infraestrutura da empresa.

Segundo o relatório, a área mais preocupante continua sendo a chamada segurança existencial. Embora diversas empresas afirmem que pretendem desenvolver sistemas de inteligência artificial geral nos próximos dois a cinco anos, o FLI conclui que nenhuma delas apresentou um plano considerado confiável para garantir o alinhamento desses modelos ou impedir eventual perda de controle sobre sistemas mais avançados.

Outro ponto destacado pelo relatório é a aproximação crescente das principais desenvolvedoras de IA com o setor de defesa. Empresas



Fonte: AI Safety Index 2026 – Future of Life Institute (FLI).

Valdo Virgo/CB/D.A Press



O que esse índice revela é que há muito a ser feito"

Arthur Igreja, especialista em tecnologia e inovação

deverão ampliar exigências como listas de fornecedores autorizados, auditoria de código, controle de tráfego e restrições ao uso de determinados modelos estrangeiros. Também devemos observar um aumento das barreiras geopolíticas entre Estados Unidos e China e regras regulatórias mais rígidas envolvendo documentação, comunicação de incidentes e avaliações de impacto”, conclui.

Já para Arthur Igreja, especialista em tecnologia e inovação e autor do livro “Conveniência é o Nome do Negócio”, o levantamento evidencia que o setor ainda está distante de atingir padrões considerados satisfatórios de segurança. “O que esse índice revela é que há muito a ser feito. Quando olhamos as dimensões e as notas, elas estão aquém. Não vemos, na verdade, nenhuma nota A. É uma tecnologia nova e, nos últimos tempos, tivemos várias notícias mostrando que a segurança e a governança dos dados passaram a ser uma prioridade”, afirma.

Segundo Igreja, a ausência de notas mais altas mostra que as empresas ainda não conseguiram cumprir os critérios estabelecidos pelo estudo. “Elas não conseguiram alcançar os standards, os benchmarks. As empresas não foram capazes de atingir esses parâmetros. Além disso, o índice permite comparar o estágio de cada companhia e mostra que algumas estão muito abaixo das líderes. A xAI, por exemplo, não está apenas um pouco abaixo da OpenAI. Ela está brutalmente abaixo”, destaca.

*Estagiário sob a supervisão de Rosana Hessel

SEBASTIEN BOZON / AFP



Anthropic é acusada por agência chinesa de deixar brechas no Claude Code

Disputa geopolítica pela liderança

A corrida tecnológica da inteligência artificial (IA) mostra várias disputas em curso nesse mercado trilionário que vai além da disputa comercial entre China e Estados Unidos. Uma das mais recentes ocorreu entre o país asiático e a norte-americana Anthropic, em torno de um mecanismo de monitoramento.

O National Vulnerability Database, órgão chinês de cibersegurança, afirmou ter identificado supostas vulnerabilidades de backdoor (porta dos fundos) no Claude Code, ferramenta de programação desenvolvida pela Anthropic. Segundo a agência, versões lançadas entre abril e junho conteriam um mecanismo capaz de transmitir informações sensíveis dos usuários, como identidade e localização geográfica, para servidores remotos sem autorização. Diante da suspeita, o órgão recomendou

que usuários removam o software ou atualizem imediatamente para a versão mais recente.

A controvérsia surgiu após discussões em fóruns on-line apontarem que a Anthropic teria inserido códigos destinados a monitorar acessos provenientes da China. Segundo a Anthropic, empresas chinesas, entre elas a Alibaba, estariam recorrendo a esse tipo de prática desde fevereiro.

O episódio elevou as tensões entre os dois países na corrida pela liderança em IA. Após a divulgação das alegações, a Alibaba proibiu oficialmente seus funcionários de utilizarem o Claude Code no ambiente corporativo.

Para o cofundador da plataforma de IA Mogno, João Mano, a disputa em torno do Claude Code reflete uma competição tecnológica

que tende a se intensificar nos próximos anos. Segundo ele, a inteligência artificial tornou-se um ativo estratégico para os países, o que dificulta qualquer tentativa de estabelecer limites globais para seu desenvolvimento. “A IA vem mostrando-se uma arma tecnológica extremamente poderosa. Qual governo ou país estaria genuinamente disposto a abrir mão dela? Em um mundo multipolar, qualquer tentativa de controle tende a fracassar diante da incerteza de que a outra parte continuaria desenvolvendo a tecnologia sem restrições”, afirma.

Na avaliação do especialista, o cenário atual lembra a corrida armamentista da Guerra Fria, mas sem mecanismos internacionais capazes de impor limites. “Na corrida nuclear, o mundo

construiu tratados, inspeções e mecanismos de dissuasão. Na inteligência artificial, sequer existe um consenso mínimo sobre o que deve ser medido, quem deve fiscalizar ou quais consequências devem existir para quem descumprir regras”, observa.

O executivo e escritor Arthur De Santos avalia que o episódio envolvendo a Anthropic também faz parte de uma disputa mais ampla pela liderança em inteligência artificial entre EUA e China. Segundo ele, existe uma guerra de narrativas, mas ambos os países possuem vantagens estratégicas distintas. “A China conta com modelos bastante maduros, um enorme ecossistema digital, grande volume de dados e uma forte coordenação entre Estado, indústria e pesquisa. Isso não pode ser subestimado”, explica. (PJ)