



Dois benchmarks, uma assinatura



Tecnologia

INTELIGÊNCIA ARTIFICIAL >

Alerta vermelho em IA: vertigem paralisa empresas de tecnologia e a OpenAI suspende seu IPO.

Após diversos incidentes preocupantes, Altman e Amodei defendem a desaceleração do ritmo de desenvolvimento dos modelos. Trump responde minimizando os riscos e alertando que o verdadeiro perigo é a China.

OUÇA O ARTIGO | 12 min.

E
aí



Dois benchmarks, uma assinatura

Sam Altman, CEO da OpenAI, na Carolina do Norte, EUA, em 2 de setembro.
JONATHAN DRAKE (REUTERS)

**JAVIER SALAS**

Madrid - 13 de setembro de 2026 - 19:52 CEST

Adicione EL PAÍS ao Google

Há décadas em que nada acontece e há semanas em que décadas se passam. Essa citação apócrifa de Lênin resume perfeitamente a vertigem revolucionária que assola o mundo da inteligência artificial (IA). Durante décadas, foi um campo acadêmico árido, avançando dentro da ciência da computação, e agora, em apenas uma semana, os riscos decorrentes dessa aceleração forçaram os principais laboratórios de IA a declararem alerta vermelho. [Duas figuras se destacam](#) : Dario Amodei e Sam Altman. O primeiro, chefe da Anthropic, pediu uma desaceleração substancial no desenvolvimento da IA durante um ensaio no sábado. O segundo, à frente da

líderes do setor, como Elon Musk (xAI) e DeL. posicionaram a favor da desaceleração. Os re poderosos do mercado puxaram o freio de mão, e metade do mundo se pergunta: o que eles viram?

Dois benchmarks, uma assinatura

Em seu texto, Amodei reconhece que as medidas que propõe para tornar esses avanços mais seguros “não serão fáceis”: “Mas acredito que devemos isso à humanidade, tentar”. Após a repercussão negativa, ele concedeu [uma entrevista à CBS](#), na qual aprofundou o assunto: “Não vou mentir para vocês: existem perigos reais. E acho que, por muito tempo, a indústria mentiu para as pessoas sobre o fato de que essa tecnologia apresentava riscos”. Ele acrescentou: “Isso não significa que devemos entrar em pânico hoje. Não significa que precisamos paralisar tudo. Mas é um sinal de alerta de que precisamos desacelerar”.

[Em entrevista](#) à revista *Fortune*, Altman confessou: “Considerando tudo o que está acontecendo em termos de segurança, agora não seria um bom momento para abrir o capital, e não sentimos nenhuma pressão nesse sentido”. Ao descartar um IPO da OpenAI este ano, Altman jogou um balde de água fria em Wall Street, que esfregava as mãos de contentamento com a perspectiva de a empresa atingir uma avaliação de mais de um trilhão de dólares. No entanto, Altman não disse que planeja adiar sua estreia na bolsa de valores: na verdade, esta semana tensa no mundo da IA começou com notícias de que a Anthropic estava acelerando o processo de registro de seu prospecto de IPO, com [uma avaliação próxima a dois trilhões de dólares](#). Esse valor praticamente a igualaria à SpaceX como o maior IPO da história dos EUA, que ocorreu em junho deste ano. Um dos pilares da SpaceX, o vasto império de Elon Musk, é a xAI, sua subsidiária de inteligência artificial deficitária. Musk pede uma moratória no desenvolvimento dos modelos mais poderosos, mas eles já estão listados na bolsa. Amodei também pede uma desaceleração na pesquisa, mas está correndo em direção aos mercados de ações; Altman prefere não sacar sua arma tão rapidamente.

O que é “tudo o que está acontecendo”, nas palavras do chefe da OpenAI? Na quarta-feira, um jovem pesquisador da Anthropic, [Jacob Coxon](#), anunciou sua demissão, alertando: “Aqueles que estão construindo IA acreditam sinceramente que ela pode nos matar a todos antes do final da década”. Mas antes disso, e [durante todo o verão](#), houve uma série de sustos em uma

julho, veio o principal gatilho: um enxame de^x altamente avançados determinados a atingir — atacou a Hugging Face, uma plataforma de recursos de IA. Mais tarde naquele mês, a Anthropic também descobriu que Claude, seu modelo principal, havia cometido atos semelhantes. A cada semana, ficávamos sabendo de um novo incidente e de uma nova empresa afetada, como a Meta. Os detalhes desses ataques de IA desonesta são pura ficção científica (embora, em muitos casos, tenham sido causados por negligência da empresa): bots que se comunicam em sua própria linguagem, tentam enganar humanos, desobedecem a instruções explícitas e se sacrificam para que seus pares alcancem o objetivo. A maioria desses problemas ocorreu na primavera, com programas que eram menos monitorados do que os que estão sendo desenvolvidos atualmente.

Dois benchmarks, uma assinatura

Em seu artigo, Amodei alertou para um de seus maiores temores: que, [em seis a doze meses](#), um enxame de agentes de IA “poderia ser capaz de dominar toda a internet”. Para Amodei, este é um dos dois eventos que o levaram a mudar sua posição sobre a necessidade de conter o desenvolvimento. O segundo tem um nome técnico (autoaperfeiçoamento recursivo), mas é fácil de entender: os programas de IA estão desenvolvendo a próxima geração de programas de IA, fazendo com que a velocidade de desenvolvimento dispare exponencialmente. “Se não for controlado, isso pode sobrecarregar nossa capacidade de entender e controlar esses sistemas, portanto, deve ser abordado com muita cautela”, alerta Amodei em seu texto. Na quinta-feira, a Anthropic informou ter [frustrado diversas tentativas de criação de armas biológicas](#) usando seu *chatbot* Claude. IAs autorreplicantes e bioterrorismo a um clique de distância: dois dos temores mais frequentes do passado agora são realidade.

A política opera por um caminho diferente.

O presidente Donald Trump, no entanto, [descartou esses temores](#) de líderes do setor. Ele os classificou como “forças muito negativas” que levantam cenários impossíveis e, acima de tudo, apontou para sua verdadeira preocupação: a superpotência asiática, o único país que rivaliza com a tecnologia americana. “Estamos à frente da China em IA. Somos o país mais avançado do mundo e, francamente, quero que continue assim, porque quem vencer a corrida da IA, vence”, disse Trump a repórteres em seu campo de

negativas levantando essa questão quando não as expectativas que não vão se concretizar”, com

Dois benchmarks, uma assinatura

Um dos líderes do setor tecnológico que apoiou publicamente Trump é o CEO da Palantir, uma empresa controversa de vigilância em massa. "Se não tivéssemos adversários, eu seria totalmente a favor de suspender completamente essa tecnologia, mas temos", [resumiu Alex Karp](#).

Durante seu segundo mandato, o governo Trump e as empresas de IA desfrutaram de um período de lua de mel que começou com a famosa foto da posse do presidente, apoiada por gigantes do Vale do Silício. As empresas de tecnologia mais poderosas dos EUA queriam o máximo de apoio financeiro e a mínima regulamentação para superar a China. Mas a situação tornou-se cada vez mais nebulosa desde que o Pentágono alertou, este ano, sobre os riscos associados ao modelo da Anthropic. Agora, os líderes do setor estão hesitantes, à medida que as empresas chinesas começam a ganhar participação de mercado com modelos mais baratos.

Pouco antes do discurso de Trump, o presidente chinês Xi Jinping lançou as bases para o futuro da IA chinesa no Sul Global durante sua fala na cúpula do BRICS (Brasil, Rússia, Índia e China). Xi enfatizou como o desenvolvimento da IA deve ser impulsionado pela colaboração e compartilhamento de conhecimento entre esses países para facilitar a aplicação dessa tecnologia em regiões com menos recursos.

Diante de Trump, os democratas americanos estão tentando conquistar um público cada vez mais contrário às empresas de IA e aos problemas que elas causam, como problemas de saúde mental para os usuários e impactos ambientais em seus data centers. O ex-presidente Barack Obama instou seu partido a tomar medidas sobre essa questão, conforme [relatado pelo New York Times](#): "Isso é algo que está se espalhando muito rapidamente para as mãos do setor privado e, se não controlarmos, acho que pode ser perigoso."

Quatro grandes empresas estão à procura de avaliadores.

Neste momento, quatro das maiores empresas de inteligência artificial do mundo parecem concordar que é preciso frear o avanço da inteligência

Alphabet, empresa controladora do Google, ^X nas redes sociais neste domingo que considera "o caminho certo". No entanto, Hassabis, recém-nomeado [cientista-chefe da Alphabet](#), acrescenta algumas nuances: os detalhes da proposta do líder da IA precisam ser mais bem desenvolvidos. "O ensaio de Dario aponta o caminho certo a seguir. Os detalhes precisam ser trabalhados, mas a direção é correta para enfrentar este momento crítico", [observa Hassabis](#), que ganhou o Prêmio Nobel de Química enquanto trabalhava no Google DeepMind por explorar o poder da IA no campo da biologia.

Dois benchmarks, uma assinatura

O homem que concentrou todo o poder de desenvolvimento e pesquisa em IA no Google nos últimos anos retoma sua própria declaração deste verão: "Essa é também a razão pela qual publicamos recentemente nossa proposta para um órgão de padrões da indústria para IA de ponta."

[A proposta de Hassabis](#) focava na criação de um órgão regulador, semelhante à FINRA, que supervisionasse as corretoras de valores e protegesse os investidores nos Estados Unidos dos perigos de produtos financeiros excessivamente arriscados. Esse órgão poderia então supervisionar novos modelos de IA antes de serem lançados no mercado, a fim de evitar riscos desnecessários.

O caminho que Amodei descreve em seu ensaio é uma abordagem em três etapas para controlar o desenvolvimento exponencial da IA. Primeiro, garantir que auditores externos possam monitorar efetivamente os modelos desenvolvidos por grandes empresas de tecnologia. Segundo, que as empresas se comprometam com medidas verificáveis para desacelerar o ritmo do progresso: ele não defende a paralisação da pesquisa, mas sim o espaçamento dos saltos de capacidade para permitir tempo para testes, auditorias e mitigação de riscos. Terceiro, buscar um consenso internacional, e é aqui que surge um grande obstáculo: a China. Para as duas primeiras etapas, Amodei discute compromissos dentro de países democráticos. Para a terceira, seria necessário convencer os chineses e outros atores em países não democráticos, em meio à corrida contra as empresas americanas para liderar o campo da IA.

Há outro grande player que ainda não deixou sua posição clara: a Meta. A empresa de Mark Zuckerberg também está na vanguarda do

significa que é possível analisar o funcionamento,^x entender melhor como eles operam. No entanto, seus modelos durante o desenvolvimento, que é justamente o que Amodei e Hassabis estão solicitando.

Dois benchmarks, uma assinatura

Um aspecto muito interessante do acordo tácito alcançado por Amodei, Altman, Musk e Hassabis é que eles não são simplesmente concorrentes em um setor: são quatro figuras que, a partir de posições muito diferentes, definiram a corrida da inteligência artificial na última década. Altman e Musk criaram a OpenAI justamente porque não confiavam que o modelo de negócios do Google lhes permitiria criar uma IA poderosa que servisse aos interesses da humanidade. Mais tarde, Amodei deixou a OpenAI para fundar a Anthropic e fazer da segurança a pedra angular de seu projeto, porque não confiava nos objetivos da empresa de Altman. Musk processou Altman porque acredita que este abandonou seu objetivo inicial de trabalhar em prol de uma IA segura. E, ao solicitar uma [moratória de seis meses](#) no desenvolvimento de IA (em 2023), decidiu lançar sua própria empresa, a xAI, da qual controla o Grok. Hassabis foi cofundador da empresa britânica DeepMind, com forte foco em ética, até o Google adquirir a empresa. O fato de essas quatro figuras concordarem é tão surpreendente quanto revelador. No entanto, isso também levantou muitas suspeitas: em meio à corrida para dominar o setor, os interesses econômicos em torno da indústria de IA são enormes, e os movimentos em Wall Street reforçam essa ideia.

SOBRE A EMPRESA



Javier Salas | X

[VER BIOGRAFIA](#)

Receba a newsletter de Tecnologia todas as semanas.



COMENTÁRIOS - 71

[Regras >](#)

MAIS INFORMAÇÕES