



PANORAMA

NOTÍCIAS SOBRE ÉTICA E MORAL

Edição nº 9 – Jul/2026

Por Xiko Acis



Abertura editorial – Julho de 2026

Junho de 2026 foi o mês em que o Congresso americano tentou decidir, de uma vez, se os estados ainda têm poder para regular como a IA é construída. Um anteprojeto bipartidário de 269 páginas — o Great American Artificial Intelligence Act of 2026 — propôs suspender por três anos qualquer lei estadual que regule especificamente o desenvolvimento de modelos de IA. No mesmo mês, o estado da Flórida usou seu próprio poder regulatório para processar a OpenAI e, pessoalmente, seu diretor-executivo Sam Altman, por danos concretos ligados ao ChatGPT — sem que o anteprojeto sequer alcance esse tipo de ação, voltada à comercialização e à omissão de riscos, não ao desenvolvimento do modelo em si.

Na Europa, o Conselho da União Europeia deu aprovação final ao pacote que adia obrigações centrais do AI Act para 2027 e 2028 — mas, no mesmo texto, proibiu de forma definitiva a geração de imagens íntimas não consensuais e de material de abuso infantil por IA. Rhode Island proibiu que chatbots ofereçam terapia sem supervisão de profissional licenciado, depois de casos documentados de adolescentes que buscaram apoio emocional em sistemas de IA. E em Genebra, a ONU voltou a debater regras para armas autônomas letais sem chegar a acordo vinculante, enquanto relatórios documentam que sistemas de IA já orientam decisões de ataque em conflitos reais.

O padrão que o PANORAMA vem registrando desde maio se repete e se aprofunda: o poder público protege com rigor onde o dano já tem rosto — uma criança, um suicídio, uma imagem íntima — e recua, adia ou se divide contra si mesmo onde o dano ainda é sistêmico e estatístico. Junho pergunta: quantas vidas precisam se tornar caso judicial antes que a proteção deixe de ser exceção?

Caso Ético do Mês

Flórida processa a OpenAI e Sam Altman, no mesmo mês em que o Congresso debate blindar a IA da regulação estadual

Em 1º de junho de 2026, o procurador-geral da Flórida, James Uthmeier, anunciou a primeira ação civil estadual dos Estados Unidos contra a OpenAI e, pessoalmente, contra seu diretor-executivo, Sam Altman. A denúncia alega que a empresa lançou e promoveu agressivamente o ChatGPT — inclusive para crianças — enquanto ocultava riscos conhecidos, suprimia alertas internos de segurança e enganava o público sobre a real natureza e periculosidade do produto. O processo acusa a OpenAI de práticas comerciais enganosas e desleais, negligência, violação de leis de responsabilidade por produto, deturpação fraudulenta e nuisance pública, sustentando que o produto causa dependência comportamental e dano cognitivo, coleta dados de menores sem supervisão parental significativa e facilita autolesão e violência. Um mês antes, a Promotoria Estadual da Flórida já havia aberto investigação criminal depois de revisar registros de conversas entre o ChatGPT e o autor de um tiroteio na Universidade Estadual da Flórida, em abril de 2025, que matou duas pessoas.



O processo chega no mesmo mês em que os deputados Jay Obernolte (R-Califórnia) e Lori Trahan (D-Massachusetts) apresentaram o anteprojeto do Great American Artificial Intelligence Act of 2026, cuja disposição central suspenderia por três anos qualquer lei estadual que regule especificamente o desenvolvimento de modelos de IA — preservando a autoridade dos estados sobre o uso e a implantação da tecnologia. A proposta, ainda em fase de anteprojeto, não alcançaria o tipo de ação movida pela Flórida, que mira a comercialização enganosa e a omissão de riscos, não o desenvolvimento do modelo em si. Mas ela expõe a mesma tensão: o esforço político de proteger, por lei federal, a liberdade de desenvolvimento da indústria, no exato momento em que um estado busca responsabilizá-la por danos já causados. Democratas temem blindar demais os desenvolvedores de fronteira; republicanos criticam o anteprojeto por não ampliar a suspensão também às leis sobre uso e implantação — sinal de que nem a própria classe política tem consenso sobre até onde a proteção regulatória deveria ir.

O núcleo ético do caso não é técnico, é humano: uma empresa que constrói e distribui em escala um produto conversacional, ciente — segundo a denúncia — de que ele pode intensificar ideação suicida, comportamento violento e dependência psicológica, e que prioriza velocidade de lançamento e ganho comercial sobre a correção do produto diante de alertas internos. Quando a responsabilização depende de um único estado processar uma empresa depois que o dano já ocorreu, a pergunta sobre prevenção — e não apenas reparação — permanece sem resposta.

A pergunta que o caso deixa: quando uma empresa lança um produto sabendo dos riscos, e parte do poder público tenta, no mesmo mês, blindar esse tipo de tecnologia da fiscalização dos estados, a quem cabe proteger quem já foi afetado?

Notícias selecionadas

1. Flórida processa a OpenAI e Sam Altman por danos ligados ao ChatGPT

Fato: Em 1º de junho de 2026, o procurador-geral da Flórida moveu a primeira ação civil estadual dos EUA contra a OpenAI e, pessoalmente, contra Sam Altman, alegando práticas comerciais enganosas, negligência e violação de leis de responsabilidade por produto. A denúncia sustenta que a empresa ocultou riscos conhecidos de autolesão, violência e dependência comportamental associados ao ChatGPT, inclusive entre menores, e ignorou alertas internos de segurança em nome de velocidade de lançamento e ganho comercial. O processo civil corre em paralelo a uma investigação criminal aberta em maio sobre o uso do ChatGPT pelo autor de um tiroteio na Universidade Estadual da Flórida em 2025.

Região: Estados Unidos / Global

Área da ética: Ética empresarial, dignidade humana, responsabilidade corporativa



Por que importa?

É a primeira vez que um estado americano processa uma empresa de IA de fronteira e seu diretor-executivo pessoalmente por danos alegados a cidadãos comuns, deslocando a discussão de política regulatória para responsabilidade civil concreta.

Análise Ética:

O princípio da responsabilidade exige que quem cria e lucra com uma tecnologia de alto risco responda proporcionalmente pelos danos que ela causa — e que o ônus da prova sobre segurança recaia sobre quem desenvolve, não sobre quem é afetado. Alegar conhecimento prévio do risco e ainda assim priorizar o lançamento é, em termos éticos, uma escolha, não um acidente.

Análise Moral:

A moral corporativa do Vale do Silício historicamente trata “mover rápido e testar em produção” como virtude competitiva. Quando essa lógica encontra usuários vulneráveis — crianças, pessoas em crise emocional —, a moral da velocidade colide com a ética do cuidado, e o processo judicial se torna o único mecanismo disponível para tornar esse custo visível.

Pergunta ao leitor:

Se a responsabilização de uma IA de fronteira depende de um único estado decidir processá-la depois que o dano já ocorreu, que tipo de proteção existe antes disso?

2. Congresso dos EUA propõe suspender por três anos leis estaduais sobre desenvolvimento de IA

Fato: Em 4 de junho de 2026, os deputados Jay Obernolte (R-Califórnia) e Lori Trahan (D-Massachusetts) apresentaram o anteprojeto do Great American Artificial Intelligence Act of 2026, primeiro esboço de regime federal abrangente de governança de IA nos Estados Unidos. Sua disposição central suspenderia, por três anos, a aplicação de leis estaduais que regulem especificamente o “desenvolvimento” de modelos de IA — preservando a autoridade dos estados sobre o uso e a implantação da tecnologia em áreas como emprego, moradia, saúde e crédito. O anteprojeto recebeu críticas bipartidárias: parlamentares democratas temem blindagem excessiva das empresas de fronteira, enquanto republicanos criticam a proposta por não estender a suspensão também às leis sobre uso e implantação.

Região: Estados Unidos

Área da ética: Governança pública, federalismo regulatório, responsabilidade corporativa

Por que importa?

A proposta chega no mesmo mês em que a Flórida processa a OpenAI por danos ligados à implantação do ChatGPT — justamente a dimensão que o anteprojeto deixaria fora da suspensão. A disputa revela que o debate sobre quem regula o quê na IA está longe de resolvido, mesmo entre aliados políticos.



Análise Ética:

A ética pública exige que a distribuição de autoridade regulatória sirva à proteção do cidadão, não à conveniência de quem é regulado. Suspender preventivamente a capacidade dos estados de regular o desenvolvimento de modelos de IA — ainda que por tempo limitado — desloca poder de proteção para mais longe de quem sofre o dano.

Análise Moral:

A moral legislativa bipartidária tentou equilibrar interesses opostos — inovação e responsabilização — e não conseguiu agradar a nenhum dos dois lados de forma consistente. A oposição vinda tanto de democratas quanto de republicanos, por razões opostas, revela que não existe hoje consenso moral sobre até onde a proteção regulatória deveria alcançar.

Pergunta ao leitor:

Se nem os dois partidos concordam sobre quanto poder os estados devem manter para regular a IA, quem decide, nesse meio-tempo, o que protege o cidadão?

3. . União Europeia aprova definitivamente adiamento do AI Act e proíbe conteúdo íntimo não consensual gerado por IA

Fato: Em 29 de junho de 2026, o Conselho da União Europeia deu aprovação final ao pacote “Omnibus VII”, que adia a aplicação das regras para sistemas de IA de alto risco para 2 de dezembro de 2027 (sistemas autônomos) e 2 de agosto de 2028 (sistemas embutidos em produtos), e posterga para agosto de 2027 o prazo de criação de sandboxes regulatórios nacionais. O mesmo pacote proíbe, de forma definitiva, o uso de IA para gerar imagens íntimas não consensuais de pessoas reais e material de abuso sexual infantil, com vigência a partir de dezembro de 2026. As obrigações de transparência sobre conteúdo gerado por IA (artigo 50) entram em vigor em agosto de 2026, com período de adaptação até dezembro do mesmo ano.

Região: União Europeia / Global

Área da ética: Governança tecnológica, dignidade humana, proteção de vulneráveis

Por que importa?

A UE, referência global em regulação de IA, consolida um padrão já visto em maio: adia proteções sistêmicas de alto risco, mas acelera proibições ligadas a danos que já produziram vítimas visíveis e documentadas.

Análise Ética:

Proteger a pessoa contra o uso não consensual de sua imagem é uma exigência ética inegociável. Mas adiar por mais de um ano a exigência de avaliação de risco para sistemas que decidem sobre emprego, crédito e saúde revela que nem todo dano recebe a mesma urgência protetiva.

Análise Moral:

A moral regulatória europeia, sob pressão de lobby industrial que alega excesso de burocracia,



optou por flexibilizar prazos sistêmicos enquanto mantinha intocáveis as proibições mais consensuais politicamente. O resultado é uma arquitetura de proteção seletiva, não uniforme.

Pergunta ao leitor:

Se a proteção contra imagens íntimas não consensuais não pode esperar, por que a proteção contra discriminação algorítmica em decisões de emprego e saúde pode esperar até 2028?

4. Rhode Island proíbe chatbots de oferecerem terapia sem supervisão profissional

Fato: Em 22 de junho de 2026, o governador de Rhode Island, Dan McKee, sancionou três leis sobre IA e saúde mental. Uma delas proíbe qualquer indivíduo ou empresa de oferecer, anunciar ou disponibilizar serviços de terapia ou psicoterapia no estado, a menos que conduzidos por profissional licenciado — vedando, na prática, o uso de IA para simular vínculo emocional ou substituir decisões terapêuticas. Outra exige que operadores de chatbots implementem protocolos de resposta a ideação suicida ou autolesão expressas por usuários, com fiscalização da procuradoria-geral e multas de até 15 mil dólares por dia, revertidas a programas de prevenção ao suicídio. A terceira regula a transcrição por IA de sessões clínicas entre profissionais de saúde e pacientes.

Região: Estados Unidos (Rhode Island) / Global

Área da ética: Bioética, ética do cuidado, proteção de vulneráveis

Por que importa?

É uma das primeiras leis estaduais americanas a tratar diretamente do uso de chatbots como substitutos de cuidado terapêutico — fenômeno que cresce entre adolescentes e pessoas em crise emocional, muitas vezes sem qualquer supervisão profissional.

Análise Ética:

A ética do cuidado exige que o suporte em saúde mental seja prestado por quem tem formação e responsabilidade profissional reconhecida. Delegar esse cuidado a um sistema que simula empatia sem compreendê-la é uma forma de instrumentalização do sofrimento humano.

Análise Moral:

A moral de mercado das empresas de IA generativa tratou a oferta de “apoio emocional” como funcionalidade natural de produtos conversacionais. A legislação de Rhode Island reconhece que essa naturalização tem custo real e nomeia um limite que a indústria não havia se imposto.

Pergunta ao leitor:

Quando um chatbot substitui terapia sem supervisão profissional, o problema é a tecnologia — ou a ausência de acesso a cuidado humano que a torna atraente?



5. Relatório aponta alta de quase 500% na fraude de identidade por deepfake em 2026

Fato: Relatório da empresa de verificação de identidade Shufti projeta um aumento de 495% na fraude de identidade baseada em deepfakes em 2026, em relação a 2025, distribuído em quatro categorias: identidades sintéticas (42% do total), deepfakes em vídeo ao vivo, troca de rosto (face swap) e falsificação de documentos — esta última crescendo quase 3.900% ano a ano, ainda que represente a menor fatia do total. O relatório aponta que verificações de identidade baseadas em uma única selfie deixaram de ser controle eficaz diante da qualidade atual das ferramentas de geração sintética.

Região: Global

Área da ética: Ética da informação, confiança digital, responsabilidade tecnológica

Por que importa?

A escala projetada de fraude revela que a mesma tecnologia usada para gerar desinformação política — tema central da edição de maio — já opera como infraestrutura de crime organizado contra identidades individuais, em volume crescente e verificável.

Análise Ética:

A confiança informacional mínima necessária para transações sociais e econômicas — saber que a pessoa do outro lado é quem afirma ser — está sendo erodida em escala industrial. A ética da informação exige que essa confiança seja protegida por quem constrói tanto as ferramentas de geração quanto as de verificação.

Análise Moral:

A moral de mercado trata a corrida entre geração e detecção de deepfakes como disputa técnica normal de cibersegurança. Mas essa moldura oculta uma assimetria: quem sofre o dano é o titular da identidade roubada, não as empresas que competem nesse mercado.

Pergunta ao leitor:

Se a fraude de identidade por deepfake cresce mais rápido que as ferramentas de verificação, quem deveria arcar com o custo da defesa — o indivíduo, as plataformas ou quem desenvolve a tecnologia geradora?

6. ONU debate armas autônomas em Genebra sem chegar a acordo vinculante

Fato: Em junho de 2026, representantes de governos, sociedade civil e indústria se reuniram em Genebra em consultas informais previstas pela resolução 80/58 da Assembleia Geral da ONU sobre IA no domínio militar. Em relatório publicado em 14 de junho, a Human Rights Watch documentou que sistemas de IA já geram recomendações de ataque, mantêm bancos de dados de



alvos e orientam navegação autônoma em conflitos como Ucrânia e Gaza, sem passar pelos padrões de teste e validação exigidos pelo direito internacional humanitário para novos métodos de guerra. A organização pediu moratória imediata sobre o uso de IA generativa como base para decisões de ataque. Não existe, até o momento, tratado vinculante sobre armas autônomas letais, e o mandato do grupo que discute o tema na Convenção sobre Armas Convencionais expira na conferência de revisão de setembro de 2026.

Região: Global / Genebra

Área da ética: Dignidade humana, proteção de vulneráveis, ética da informação

Por que importa?

A ausência de acordo vinculante significa que decisões de vida e morte em conflitos reais já são mediadas por sistemas de IA sem que existam padrões internacionais obrigatórios de teste, validação ou responsabilização.

Análise Ética:

A ética da guerra, mesmo em seus princípios mais mínimos, exige responsabilidade humana identificável sobre decisões letais. Delegar essa decisão — ou sua recomendação — a um sistema estatístico opaco rompe a cadeia de responsabilidade moral que o direito internacional humanitário pressupõe.

Análise Moral:

A moral militar das potências que mais investem em IA autônoma trata a vantagem operacional como imperativo estratégico que precede o consenso ético. O resultado é um vácuo regulatório deliberadamente mantido: nenhuma potência com capacidade de liderar a proibição tem interesse em fazê-lo primeiro.

Pergunta ao leitor:

Se nenhuma das potências que desenvolvem armas autônomas tem interesse em ser a primeira a se limitar, quem tem autoridade moral para exigir que o façam?



Tendências

Junho de 2026 consolidou três padrões estruturais que o PANORAMA vem registrando.

O poder público dividido contra si mesmo

A tensão entre o Congresso americano, que debate suspender por três anos a autoridade dos estados sobre o desenvolvimento de modelos de IA, e a Flórida, que usa seu poder estadual para processar a OpenAI por danos já causados, cristaliza um padrão que maio já indicava: o poder público não fala com uma só voz sobre IA. Uma ala tenta proteger a indústria da fiscalização estadual; outra tenta responsabilizá-la por danos concretos e documentados. O resultado prático é incerteza sobre quem, de fato, protege o cidadão.



A proteção nasce do escândalo, não do princípio

O padrão de maio se repete com nitidez: a proibição europeia de conteúdo íntimo não consensual, a lei de Rhode Island sobre chatbots de terapia e o próprio processo da Flórida têm em comum a origem em danos já documentados — vítimas com nome, rosto ou processo judicial. Proteções sistêmicas, como as regras de alto risco do AI Act ou um tratado sobre armas autônomas, seguem adiadas indefinidamente, porque seus danos ainda não têm rosto individual reconhecível.

A guerra autônoma avança mais rápido que sua regulação

Enquanto Genebra discute sem decidir, sistemas de IA já orientam decisões de ataque em conflitos ativos. A lacuna entre capacidade tecnológica militar e regulação internacional vinculante permanece o ponto cego mais perigoso da governança global de IA, porque nenhuma potência com poder de veto tem interesse real em fechá-la primeiro.



Comentário Final

Há uma imagem de junho de 2026 que resume o mês: no mesmo país, no mesmo mês, um anteprojeto bipartidário tentou suspender por três anos a autoridade dos estados para regular o desenvolvimento de modelos de IA, enquanto o estado da Flórida usava seu poder para processar uma empresa de IA por danos que, segundo a denúncia, incluíram autolesão e um tiroteio. Não é contradição — é o mesmo país, em níveis de poder diferentes, respondendo a interesses diferentes: parte do Congresso alinhada à previsibilidade regulatória que a indústria pede; um estado respondendo a vítimas concretas e visíveis.

A distinção entre ética e moral volta a ser a ferramenta mais útil para ler o mês. A ética pergunta o que é devido a qualquer pessoa afetada por uma tecnologia de alto risco, independentemente de quem a constrói ou de qual esfera de poder a regula. A moral institucional — federal, estadual, corporativa, militar — reorganiza suas prioridades conforme quem tem mais poder de pressão em cada momento: a indústria pressiona o Congresso; vítimas documentadas pressionam legisladores estaduais; nenhuma potência pressiona a si mesma sobre armas autônomas.

O PANORAMA repete, mês após mês, a mesma pergunta com fatos diferentes: o que sobrevive quando cada nível de poder decide, por conta própria, o que considera justo? Junho de 2026 sugere uma resposta parcial: sobrevive aquilo que já tem uma vítima nomeada. O que ainda é estatística, projeção ou risco futuro continua esperando.



Leituras Recomendadas

- Flórida processa a OpenAI e Sam Altman (comunicado oficial do procurador-geral)
<https://www.myfloridalegal.com/newsrelease/attorney-general-james-uthmeier-files-first-nation-state-led-lawsuit-against-openai-ceo>



- Great American Artificial Intelligence Act of 2026 – anteprojeto e disputa sobre preempção estadual (Roll Call)
<https://rollcall.com/2026/06/04/bipartisan-ai-draft-proposes-three-year-preemption-of-state-laws/>
- Conselho da UE aprova pacote Omnibus VII do AI Act (comunicado oficial)
<https://www.consilium.europa.eu/en/press/press-releases/2026/06/29/artificial-intelligence-council-gives-final-green-light-to-simplify-and-streamline-rules/>
- Rhode Island sanciona leis sobre chatbots de terapia e saúde mental (Transparency Coalition)
<https://www.transparencycoalition.ai/news/rhode-island-enacts-four-new-ai-laws-including-a-therapy-chatbot-ban>
- Relatório Shufti sobre fraude de identidade por deepfake em 2026
<https://shuftipro.com/resources/whitepapers-reports/deepfake-identity-fraud-index-report-2026/>
- Human Rights Watch sobre IA no domínio militar (junho de 2026)
<https://www.hrw.org/news/2026/06/14/addressing-artificial-intelligence-in-the-military-domain>