

TECNOLOGIA

Meta processa brasileiros por uso de deepfakes em anúncios falsos de produtos de saúde

Foram acionadas judicialmente pessoas suspeitas de usar imagens e vozes alteradas de celebridades para promover produtos de saúde falsos

A Meta informou que está processando pessoas e empresas do Brasil, da China e do Vietnã que usavam a imagem de celebridades e marcas conhecidas para veicular anúncios falsos nas plataformas da empresa.

No Brasil, os envolvidos citados pela big tech são acusados de usar deepfakes de personalidades famosas para vender produtos de saúde falsos. A empresa afirmou que os anúncios fraudulentos são feitos para parecer reais e nem sempre é fácil detectá-los. “Esse esquema, conhecido como ‘isca de celebridade’, prejudica a confiança das pessoas e viola nossas políticas”, afirma a companhia.

Segundo comunicado publicado no site oficial da dona do Facebook e do Instagram, foram acionadas judicialmente pessoas suspeitas de usar imagens e vozes alteradas de celebridades para promover produtos de saúde falsos. Também foram processadas empresas de suplementos e de treinamento que, segundo a companhia, “fazem parte de uma operação que usa deepfakes de um médico proeminente para anunciar produtos de saúde sem aprovação regulatória”.

O grupo também é acusado de vender cursos para ensinar as táticas de falsificação.

Entre as personalidades vítimas da manipulação está o médico Drauzio Varella, conforme relato publicado em sua coluna na Folha em outubro de 2025. À época,



o médico afirmou que tentava há anos - sem sucesso - retirar das redes sociais os vídeos falsos que utilizam sua imagem, e que chegou a levar o caso ao Ministério Público de São Paulo.

“Claro que fico revoltado quando vejo meu nome achincalhado por gente da pior espécie em conluio com as plataformas, depois de quase 60 anos de profissão. É assustador ver pessoas esclarecidas caírem nessas armadilhas, por acreditar que estou indicando produtos capazes de curar diabetes, dores nas costas, neuropatias periféricas e emagrecer 20 quilos em um mês”, escreveu o médico.

Ele também descreveu a dificuldade de entrar em contato com a big tech. “Com a minha insistência a cada vídeo novo, meus emails simplesmente deixaram de ser respondidos. No máximo, vinha uma resposta automática dizendo que a

publicação respeitava as normas da plataforma”, contou.

A big tech afirma que houve suspensão de métodos de pagamento, desativação de contas e bloqueio dos nomes de domínio de sites vinculados aos envolvidos. O comunicado acrescenta que “as ações judiciais e os esforços contínuos para combater golpes enviam uma mensagem clara: aqueles que buscam explorar outras pessoas em nossas plataformas serão responsabilizados”.

Personalidades em todo o mundo sofrem com o problema. O comentarista-chefe de economia do Financial Times, Martin Wolf, também já denunciou o problema em coluna publicada pela Folha. Ele relatou que vídeos com deepfakes de seu rosto convidando pessoas para um grupo de investimentos falso alcançaram pelo menos 970 mil usuários das plataformas da Meta apenas na UE (União Europeia).

Processos em outros países

Na China, a Meta está processando a empresa Shenzhen Yunzheng Technology Co. pelo uso de imagens de celebridades para anúncios falsos que atingiram usuários nos Estados Unidos, Japão e outros países. A Meta afirma que essas ações fazem parte de “um esquema de fraude maior que atraía pessoas para participar de supostos grupos de investimento”.

Também está sendo processada a empresa vietnamita Lý Van Lâm. Segundo a Meta, o grupo usou métodos de camuflagem para burlar o processo de revisão de anúncios da plataforma e veiculou anúncios falsos oferecendo grandes descontos de marcas conhecidas, como a grife Longchamp. Os usuários eram redirecionados para sites onde informavam da-

dos de cartão de crédito para comprar itens que nunca receberiam e, posteriormente, recebiam cobranças recorrentes não autorizadas.

O comunicado da empresa também afirma que, em uma outra operação em colaboração com autoridades policiais do Reino Unido e da Nigéria, foram efetuadas sete prisões vinculadas a uma central de golpes.



JOSÉ MILAGRE
Facebook.com/drjosemilagre
Instagram.com/drjosemilagre
Youtube.com/josemilagre

IA na saúde: o que muda com a nova Resolução do CFM

O uso da inteligência artificial na saúde deixou de ser tendência para se tornar parte do cotidiano de hospitais, clínicas e consultórios. Sistemas auxiliam análises, apoiam decisões e organizam processos. Em 2026, a questão não é mais se a IA será utilizada, mas como adotá-la sem comprometer segurança jurídica, ética profissional e proteção do paciente. A nova Resolução nº 2.454/2026 do Conselho Federal de Medicina oferece respostas relevantes e inaugura um novo marco regulatório.

O limite ético da autonomia profissional

A nova norma estabelece um ponto essencial: o julgamento clínico não pode ser delegado a algoritmos. Ferramentas tecnológicas podem identificar padrões e sugerir caminhos, mas não compreendem contexto, sofrimento, histórico subjetivo ou nuances de comportamento. A decisão deve permanecer humana, pois é baseada em responsabilidade profissional, análise crítica e avaliação individualizada. Adotar recomendações automatizadas sem reflexão não representa autonomia, mas abdicação de responsabilidade.

Transparência e proteção de dados no centro do cuidado

O paciente tem direito de ser informado sempre que a IA participar do seu atendimento e pode recusar seu uso. Em um cenário em que dados pessoais se tornaram sensíveis e valiosos, transparência deixou de ser formalidade e passou a ser condição de validade. A resolução reafirma o dever de observância à Lei Geral de Proteção de Dados (LGPD), reforçando que o uso de IA na saúde envolve dados sensíveis e exige governança compatível com o nível de risco.

Governança interna como requisito obrigatório

Um dos pontos mais inovadores é a criação de Comissões de Inteligência Artificial e Telemedicina em instituições que utilizarem sistemas próprios. Isso marca a transição da IA do campo experimental para o campo da infraestrutura crítica. As comissões serão responsáveis por auditorias, revisão de modelos, análise de vies, monitoramento de desempenho e elaboração de relatórios técnicos. A IA passa a ser tratada como processo permanente, e não como ferramenta ocasional.

Risco não é destino, é gestão contínua

A resolução introduz uma classificação que diferencia riscos baixos, moderados, altos e inaceitáveis. Essa escala orienta o nível de vigilância esperado. Soluções administrativas exigem menor controle, enquanto ferramentas que influenciam decisões clínicas críticas exigem monitoramento rigoroso e constante. A mensagem é clara: quanto maior o impacto da tecnologia no cuidado, maior a responsabilidade que recai sobre as instituições e sobre o profissional.

Humano demais para ser substituído

O documento reforça que a tecnologia pode ampliar o alcance da técnica, mas não substitui componentes essenciais da prática clínica. A IA não interpreta sofrimento, não percebe hesitações e não estabelece vínculo. Cuidar exige humanidade, julgamento e sensibilidade. O país entra em uma fase na qual a inteligência pode ser artificial, mas o cuidado continua profundamente humano.

José Milagre
Colunista de tecnologia & inovação do JC, especialista em direito e tecnologia, sociedade e segurança digital, perito em informática, diretor do Instituto de Defesa do Cidadão na Internet (IDCI), Mestre e doutorando pela Unesp. Escreve aos domingos no JC.

Envie suas dúvidas, eventos e iniciativas na área de tecnologia, segurança, startups e inovação e comentários para consultor@josemilagre.com.br