

LIMITES ÉTICO-JURÍDICOS DO USO DA INTELIGÊNCIA ARTIFICIAL NA PRÁTICA MÉDICA: UMA ANÁLISE BIOÉTICA E NORMATIVA

Eliezer Vofchuk¹

Maurício Paes Manso²

RESUMO

Examina os limites éticos e jurídicos da aplicação da inteligência artificial (IA) na prática médica, que une eficiência tecnológica sem agredir direitos fundamentais. Trata-se de estudo teórico-normativo, com base em pesquisa bibliográfica e documental, que articula fundamentos da bioética (autonomia, beneficência, não maleficência e justiça) ao ordenamento brasileiro (CF/88, LGPD, CDC e Código de Ética Médica), além de análise comparada entre o AI Act da União Europeia e o modelo regulatório do FDA nos Estados Unidos. Identifica riscos clínicos e jurídicos decorrentes de opacidade decisória, vieses algorítmicos e automação bias, bem como lacunas na distribuição de responsabilidades entre médico, instituição e desenvolvedor. Demonstra que o uso de IA em saúde deve permanecer como apoio à decisão clínica, sob supervisão humana obrigatória, transparência e rastreabilidade. Propõe diretrizes de conformidade: governança e auditoria algorítmica, gestão de vieses, direito de informação e de revisão humana, registros de uso e responsabilidade proporcional e compartilhada. Conclui ser imprescindível marco regulatório setorial no Brasil que assegure explicabilidade, proteção de dados sensíveis e segurança do paciente, preservando a dignidade humana e a autonomia na relação médico-paciente.

Palavras-chave: Inteligência artificial em saúde; Bioética; Responsabilidade civil médica; Proteção de dados; Regulamentação.

INTRODUÇÃO

A inteligência artificial (IA) representa uma das mais significativas transformações tecnológicas do século XXI, configurando-se como um campo da ciência da computação voltado ao desenvolvimento de sistemas capazes de simular a cognição humana aprendendo, raciocinando e tomando decisões de modo autônomo. A inteligência artificial busca desenvolver sistemas capazes de simular capacidades humanas, como raciocínio, aprendizado, planejamento e criatividade. O principal objetivo é criar mecanismos que percebam o ambiente em que estão inseridos, processar informações e solucionar problemas de maneira autônoma, adaptando o seu comportamento com base em dados e experiências anteriores. Na área da saúde, a inteligência artificial tem contribuindo de maneira relevante na modernização da prática médica. Integrando diferentes tipos de dados, com informações genéticas, clínicas, e comportamentais, esses sistemas auxiliam os profissionais de saúde a formular diagnósticos com maior precisão.

A aplicação da IA na medicina cresce de forma exponencial e multifacetada. Na radiologia, algoritmos auxiliam na detecção de fraturas, tumores e hemorragias intracranianas com precisão comparável à de especialistas humanos. Na oncologia, ferramentas de aprendizado profundo analisam lâminas histopatológicas e identificam padrões de malignidade invisíveis ao olho humano. A cardiologia utiliza a IA na interpretação de eletrocardiogramas e no monitoramento remoto de pacientes com risco de arritmia ou insuficiência cardíaca. Já na dermatologia e oftalmologia, sistemas inteligentes realizam triagem de le-

¹ Graduando em Direito da Universidade Santo Amaro, SP.

² Orientador, Universidade Santo Amaro, SP.

sões cutâneas e detectam retinopatia diabética por meio de imagens digitais. Em unidades de terapia intensiva e emergências hospitalares, algoritmos preditivos antecipam quadros de sepse e auxiliam na priorização de pacientes em estado crítico. Essas inovações demonstram o potencial da IA em elevar a eficiência e a precisão diagnóstica, contribuindo para uma medicina mais proativa e preventiva. Atualmente, sistemas de computador aplicados à saúde já estão inseridos a diferentes especialidades médicas, demonstrando grande potencial de inovação. Na oftalmologia, programas reconhecidos por agências internacionais como a Food and Drug Administration (FDA) fazem triagem de retinopatia diabética sem necessidade de intervenção direta do médico na fase inicial. Há ainda aplicações significativas na triagem hospitalar, sobretudo em serviços de urgência, em que ferramentas auxiliam a priorizar pacientes com base no risco clínico e prever quadros graves como sepse horas antes de apresentarem sintomas críticos. Assim, essas tecnologias ampliam a capacidade diagnóstica, aceleram decisões clínicas e podem contribuir para reduzir filas de atendimento e desigualdades de acesso.

Contudo, a utilização dessas ferramentas no cotidiano médico impõe novos questionamentos e desafios. Uma das principais questões éticas está relacionado à transparência na decisão dos algoritmos. Grande parte dos modelos atuais são baseados em redes neurais, em que a lógica de funcionamento é incerta até mesmo para os desenvolvedores, tornando verdadeiras caixas-pretas. Essa característica influencia diretamente o princípio da autonomia do paciente, tendo em vista que a ausência de explicação inviabiliza o consentimento informado, violando-se o princípio da beneficência quando decisões automatizadas são tomadas sem critérios clínicos compreensíveis. Além disso, a constante delegação das decisões médicas a sistemas de IA ameaça a relação médico-paciente, que sempre foi fundada na confiança, diálogo e responsabilidade profissional. Substituir o juízo clínico exclusivamente humano por um modelo automatizado sem supervisão adequada acaba transferindo o cuidado em saúde a um agente que não tem aprendizado e discernimento ético, incapaz de responsabilidade moral.

No campo jurídico, os desafios são também expressivos. O uso de sistemas de decisão automatizada em saúde envolve tratamento de dados clínicos e genéticos, classificados como dados sensíveis pela legislação brasileira. A Lei Geral de Proteção de Dados (LGPD) estabelece princípios rígidos para o tratamento dessas informações, impondo bases legais, como limites para compartilhamento e o dever de transparência. Entretanto, muitos sistemas empregados em hospitais e clínicas operam com falta de clareza sobre a governança de dados, fazendo com que os pacientes fiquem expostos a riscos de violação de privacidade, discriminação algorítmica e mercantilização indevida de informações pessoais por empresas de tecnologia. Além disso, o ordenamento jurídico brasileiro não possui regulamentação específica sobre responsabilidade civil em caso de danos decorrentes de decisões algorítmicas. Existem incertezas quanto à imputação de responsabilidade: deve-se atribuir o dano ao profissional de saúde que utilizou o sistema, ao hospital que o implementou, ao desenvolvedor do software ou à empresa fornecedora da tecnologia?

O Conselho Federal de Medicina (CFM) reconhece oficialmente a crescente utilização de sistemas computacionais no exercício profissional, mas reafirma que a responsabilidade pela decisão final deve permanecer humana, sendo vedada a substituição do médico por sistemas automatizados. De acordo com a resolução nº 2.314/2022, tais recursos devem atuar como instrumentos de apoio e jamais como agentes autônomos de decisão clínica. Contudo, a exigência de supervisão humana por si só não é suficiente para eliminar os riscos de dependência tecnológica, erros de sistema decorrentes de vieses nos dados que alimentam os modelos. O atual cenário é de grande tensão entre inovação tecnológica e a garantia de direitos fundamentais, principalmente o direito à dignidade humana, à saúde e à proteção de dados pessoais.

É evidente que o debate em torno da aplicação de sistemas inteligentes na medicina não pode restringir-se ao entusiasmo tecnológico. É imprescindível estabelecer limites éticos e marcos jurídicos que assegurem o uso responsável dessas tecnologias, equilibrando a eficiência técnica com a responsabilidade so-

cial. Considerando que essas ferramentas já fazem parte da prática médica e continuarão a evoluir, é imprescindível regulamentar sua atuação de forma que contribua, e não substitua, a racionalidade humana e os princípios bioéticos. Dessa forma, este trabalho tem como objetivo analisar os desafios éticos e jurídicos decorrentes do uso de tecnologias algoritmizadas na medicina, demonstrando que sua adoção deve ocorrer com transparência, e supervisão humana, segurança de dados e responsabilidade profissional definida.

Diante do cenário exposto, este trabalho tem como objetivo analisar os limites éticos e jurídicos do uso da IA na prática médica, destacando a importância de sua implementação de maneira responsável, com transparência e supervisão por profissionais da saúde. Busca-se demonstrar que o avanço tecnológico, é inevitável e benéfico, devendo sempre estar subordinado à preservação da dignidade humana, à proteção de dados sensíveis e à manutenção da autonomia profissional e do paciente. Assim, mais do que uma ferramenta de eficiência, a IA deve ser compreendida como um instrumento a serviço da ética, da justiça e da vida.

FUNDAMENTOS BIOÉTICOS APLICADOS

presente trabalho busca responder à seguinte pergunta-problema: De que modo a falta de normatização do uso da IA na prática médica compromete a dignidade da pessoa humana, a autodeterminação dos pacientes e a segurança jurídica dos profissionais de saúde, e quais parâmetros ético-jurídicos poderiam ser adotados para diminuir esses riscos?

A bioética surge como um campo interdisciplinar voltado à reflexão sobre as implicações morais das práticas biomédicas e das inovações tecnológicas que afetam os seres humanos. Teve origem por volta da década de 1970, com a proposta de Van Rensselaer Potter, que concebeu a bioética como uma “ponte para o futuro”, isto é, um espaço de integração entre a biologia e os valores humanos, voltado à preservação da vida e da dignidade (Potter, 1971). Desde então, consolidou-se como disciplina fundamental para o exame ético nas decisões em saúde, principalmente em con-

textos marcados pela complexidade técnica e pela incerteza científica.

No cenário atual, a aplicação da IA à medicina intensifica a necessidade de uma reflexão bioética atualizada. A utilização de algoritmos com capacidade de processar grandes volumes de dados clínicos e de auxiliar no diagnóstico ou tratamento de pacientes apresentam novas maneiras de poder e de responsabilidade. Conforme Beauchamp e Childress (2013), a prática biomédica deve-se pautar por quatro princípios fundamentais: autonomia, beneficência, não maleficência e justiça. Esses princípios, embora utilizados para orientar a relação entre profissionais de saúde e pacientes, são também aplicáveis às tecnologias digitais em saúde, na medida em que a IA também participa do processo decisório e influencia diretamente o cuidado médico.

O princípio da autonomia exige que o paciente seja informado e possa decidir livremente sobre procedimentos diagnósticos ou terapêuticos. Todavia, a complexidade dos sistemas de IA torna muito difícil a compreensão de como os algoritmos chegam a determinadas conclusões, conhecido como opacidade algorítmica (Floridi, 2018). Essa falta de transparência desafia a efetividade do consentimento informado, uma vez que o paciente não tem informações suficientes para avaliar os riscos e benefícios. Esse cenário demanda o fortalecimento da ética da explicabilidade e da rastreabilidade, de modo que a confiança nas decisões médicas mediadas por IA seja também preservada.

Já os princípios da beneficência e da não maleficência tem obrigação de promover o bem-estar do paciente e principalmente, de evitar danos. No caso da IA, esses princípios trazem uma nova dimensão, pois o dano pode decorrer não apenas da conduta do médico, mas de erros sistêmicos ou de vieses incorporados ao treinamento do algoritmo. Mittelstadt et al. (2016) observam que os sistemas de aprendizado de máquina podem reproduzir ou amplificar desigualdades preexistentes nos dados, gerando injustiças diagnósticas e terapêuticas. Assim, a beneficência e a justiça exigem não apenas a boa intenção do agente humano, mas também a avaliação crítica das bases de dados e dos critérios técnicos que orientam as decisões automatizadas.

O princípio da justiça, impõe a distribuição equitativa dos benefícios e encargos decorrentes das práticas médicas. Aplicado à IA, esse princípio exige que o acesso às tecnologias de ponta não seja destinado somente a grupos privilegiados e que sua adoção não aprofunde ainda mais as desigualdades regionais ou socioeconômicas em saúde. A Organização Mundial da Saúde (OMS) (WHO, 2021) defende que a equidade deve ser um eixo central na governança ética da IA, o que inclui garantir diversidade nos dados utilizados e participação social na definição de prioridades tecnológicas.

A reflexão bioética também se expande com a contribuição de Hans Jonas (2006), que propõe o “princípio responsabilidade” como fundamento moral da civilização tecnológica. Para o autor, toda ação técnica deve considerar suas possíveis consequências futuras sobre a vida humana e o meio ambiente. Aplicado à medicina digital, esse princípio orienta o dever de precaução no uso da IA, especialmente em contextos nos quais erros algorítmicos podem comprometer vidas humanas. Trata-se de uma ética voltada ao futuro, que reconhece o poder transformador da tecnologia e impõe prudência ao seu uso.

Sob a ótica jurídica, os fundamentos bioéticos convergem com os princípios constitucionais da dignidade da pessoa humana, da autonomia da vontade e da proteção à vida e à saúde, previstos nos artigos 1º, III, e 196 da Constituição Federal de 1988. A bioética, nesse sentido, não se opõe ao Direito, mas o complementa, fornecendo critérios morais que auxiliam a interpretação das normas jurídicas diante de novas realidades tecnocientíficas. O direito à autodeterminação informativa, conforme sustenta Doneda (2006), constitui expressão contemporânea da dignidade humana e reforça a necessidade de consentimento e transparência no tratamento de dados pessoais em saúde. O vínculo entre bioética e responsabilidade civil emerge, portanto, como elemento essencial para a proteção do paciente em ambientes clínicos mediados por tecnologias inteligentes.

Por fim, a aplicação da IA na saúde demanda uma ética da corresponsabilidade. Médicos, programadores, gestores hospitalares e desenvolvedores de tecnologia compartilham de-

veres de diligência de forma prudente. A bioética aplicada à IA não se limita à análise das intenções individuais, mas exige uma governança ética coletiva, orientada pela proteção integral do ser humano. Assim, compreender os fundamentos bioéticos é condição fundamental para determinar os limites ético-jurídicos da IA na prática médica, bem como assegurar que o avanço tecnológico continue a servir ao propósito central da medicina: o cuidado responsável e humanizado.

Os fundamentos bioéticos que orientam o uso ético da tecnologia médica foram sistematizados de forma clássica por Tom L. Beauchamp e James F. Childress, em sua obra “Principles of Biomedical Ethics”. Segundo os autores, os quatro princípios principais autonomia, beneficência, não maleficência e justiça devem guiar as condutas clínicas e políticas públicas de saúde. O respeito à autonomia requer que os pacientes tenham a capacidade de entender e consentir livremente sobre as decisões clínicas que os afetam. A beneficência exige que os profissionais ajam para promover o bem-estar dos pacientes, enquanto a não maleficência impõe o dever de evitar danos. Já o princípio da justiça está relacionado à equidade na distribuição dos recursos e no acesso às inovações terapêuticas, especialmente diante das desigualdades estruturais do sistema de saúde brasileiro. O uso da IA, portanto, deve ser examinado à luz desses princípios para que se garanta uma prática médica ética e centrada no ser humano. O exame da literatura e das experiências internacionais mostra que a IA tem um potencial enorme como ferramenta de apoio à decisão clínica em contextos de alto risco. Seu uso pode reduzir erros diagnósticos, aumentar a eficiência terapêutica e fornecer subsídios valiosos em situações de emergência. No entanto, a ausência de uma regulação específica, aliada a limitações técnicas e éticas, impõe cautela.

A proteção dos dados pessoais configura-se como um elemento estruturante do debate sobre IA em saúde, especialmente porque o tratamento em massa de informações médicas envolve questões éticas, jurídicas e existenciais. Doneda (2011, p. 89) afirma que a autodeterminação informativa é pressuposto da dignidade humana, na medida em que assegura ao indivíduo o controle sobre o fluxo de

informações que lhe dizem respeito. Dessa forma, considerando que a IA opera a partir de grandes bases de dados sensíveis, a sua utilização na saúde só pode ser juridicamente legítima quando respeitada a privacidade do paciente, garantindo transparência, finalidade e proteção contra usos abusivos de dados.

Assim, a fundamentação teórica revela que o uso da IA em casos que envolvem alto risco médico não pode ser analisado apenas sob a ótica da eficiência técnica. É necessário integrá-lo aos paradigmas da bioética e do direito, para que o avanço tecnológico se dê sem prejuízo da dignidade humana, da autonomia dos pacientes e da segurança dos profissionais de saúde. Doneda (2011, p. 89) observa que a autodeterminação informativa é pressuposto essencial da dignidade humana, pois garante ao indivíduo controle sobre os dados pessoais que o identificam. A utilização da IA em saúde, dependente de dados sensíveis, precisa observar rigorosamente essa garantia.

RESPONSABILIDADE CIVIL NA PRÁTICA CLÍNICA COM INTELIGÊNCIA ARTIFICIAL

O modelo clássico de responsabilidade civil médica baseia-se na culpa, ou seja, na verificação de negligência, imprudência ou imperícia por parte do especialista. Porém, com a introdução da IA no processo decisório, essa lógica tradicional é relativizada. Surge a possibilidade de decisões clínicas serem parcialmente atribuídas a sistemas computacionais gerando dificuldades na determinação do nexo causal e na imputação da conduta lesiva. A doutrina já debate a aplicação de uma responsabilidade solidária ou até mesmo objetiva, especialmente nos casos em que o especialista segue a recomendação de um sistema validado institucionalmente. Há quem defenda a criação de um regime especial de responsabilidade digital na saúde, que envolva o hospital, o desenvolvedor do sistema, o especialista de saúde e até o Estado como responsável solidário, conforme o artigo 37, §6º da Constituição Federal. O tema carece de jurisprudência consolidada, mas decisões recentes sobre erro médico já começam a considerar o uso de

ferramentas tecnológicas como elemento agravante ou atenuante da responsabilidade. A responsabilidade civil tem por finalidade reparar danos causados a outrem. Nos termos do artigo 927 do Código Civil, “aquele que, por ato ilícito, causar dano a outrem, fica obrigado a repará-lo”. A responsabilidade médica tradicionalmente é tratada como responsabilidade subjetiva, que exige comprovação de culpa (negligência, imprudência ou imperícia), conforme artigo 186 do Código Civil.

Para Maria Helena Diniz (2015, p. 47), a responsabilidade médica decorre do “dever de cuidado na condução técnica da atividade profissional, cujo descumprimento configura ato ilícito reparável”.

A responsabilidade civil extracontratual do programador é regida pelo artigo 798º do Código Civil Português, que estabelece que: “Aquele que, sem culpa, causar danos a outrem, fica obrigado a repará-los, se não provar que houve culpa do lesado ou que o dano foi causado por força maior ou por caso fortuito.”

O programador também pode ser responsabilizado criminalmente por danos causados por um software. As leis penais aplicáveis incluem: Código Penal Português: O artigo 15º do Código Penal Português define o crime como a violação de um bem jurídico protegido pela lei penal, com dolo ou culpa. A Lei de Informática, Lei nº 41/2004, de 18 de agosto, estas sanções para atos ilícitos praticados por meio de informática.

Os tribunais têm sido cada vez mais exigentes com os programadores, exigindo que eles tomem todas as medidas necessárias para evitar danos causados por seus softwares.

A introdução de sistemas de IA no contexto médico-clínico adiciona uma nova camada de complexidade ao regime jurídico da responsabilidade civil na saúde. Embora a IA se apresente como ferramenta tecnológica de apoio à decisão clínica, capaz de aumentar a precisão diagnóstica e reduzir erros humanos, sua utilização também cria espaço de incerteza normativa, sobretudo quanto à atribuição de culpa, ao dever de informação e à segurança do paciente. Isso ocorre porque a IA atua como um elemento intermediário entre a conduta do profissional de saúde e o resultado final da

ação médica, constituindo o que a doutrina tem descrito como risco tecnológico assistido (Floridi, 2018).

A doutrina civilista reconhece que o padrão de diligência profissional é avaliado à luz da técnica disponível e do estado da arte (Cavaliere Filho, 2019). Assim, quando o médico utiliza sistemas de IA em sua atuação profissional, o dever de cuidado se amplia, de maneira em que se deve avaliar criticamente as recomendações algorítmicas, sem atribuir a essa tecnologia um caráter absoluto. A literatura técnica reconhece que sistemas de IA podem produzir erros algorítmicos decorrentes de vieses de dados, falhas de treinamento ou limitações estatísticas (Mittelstadt et al., 2016). Portanto, a confiança cega na máquina configura conduta negligente, pois viola o princípio ético da prudência clínica e o dever jurídico de diligência profissional.

O Código de Ética Médica brasileiro (Resolução CFM nº 2.217/2018) reforça esse entendimento ao estabelecer, em seu art. 1º, que o alvo da medicina é a saúde do ser humano, devendo o médico agir com o melhor de sua capacidade profissional e moral.

Além disso, o art. 32 veda ao médico delegar a outros profissionais atos ou atribuições exclusivos da profissão médica, mesmo que sob sua supervisão, salvo em caso de urgência ou emergência, no âmbito e nos limites da lei", o que inclui a transferência indevida da decisão clínica a sistemas de IA.

O dever de informação, decorrente do princípio da autonomia do paciente, está diretamente implicado no uso da IA. Conforme Beauchamp e Childress (2013), não há autonomia sem informação adequada e transparente. Tal entendimento encontra fundamento jurídico no art. 6º, III, do Código de Defesa do Consumidor, que garante ao paciente o direito à informação clara e adequada sobre riscos e benefícios do tratamento. No contexto de IA, isso implica informar ao paciente que parte do diagnóstico ou recomendação terapêutica será realizada com suporte algorítmico, o que inclui comunicar eventuais limitações ou riscos do sistema utilizado. Esse dever também se conecta com a LGPD (Lei nº 13.709/2018), que exige transparência no tratamento de dados pessoais sensíveis (art. 6º), categoria que inclui dados de saúde.

A problemática se intensifica quando se observa que a IA utiliza dados pessoais em massa na construção de deduções clínicas. Doneda (2006) alerta que qualquer tratamento informacional pode representar risco à privacidade e dignidade humana se sem garantias. Assim, o uso de IA sem uma base ética ou jurídica viola o direito fundamental à autodeterminação informativa. Ademais, o art. 20 da LGPD dispõe que "O titular dos dados tem direito a solicitar a revisão de decisões tomadas unicamente com base em tratamento automatizado de dados pessoais que afetem seus interesses.

Casos internacionais demonstram a relevância do tema. Em 2018, o sistema IBM Watson for Oncology sugeriu condutas terapêuticas contraindicadas para pacientes reais com câncer, chegando inclusive a recomendar tratamentos errados e até letais em simulações clínicas (Topol, 2019). Em 2019, o hospital London Royal Free Hospital foi investigado no Reino Unido por uso indevido de dados sensíveis de pacientes no treinamento do algoritmo DeepMind Health, do Google, sem consentimento informado (Floridi, 2018). Esses casos demonstram que a IA pode falhar não apenas tecnicamente, mas eticamente e juridicamente, gerando violações de privacidade e autonomia.

Em resumo, a responsabilidade civil na prática clínica com IA deve seguir três diretrizes fundamentais: (i) o médico continua sendo o responsável final pelo ato clínico, com dever reforçado de cautela; (ii) a IA deve ser utilizada como ferramenta de apoio e nunca como substituto do julgamento médico; e (iii) a responsabilidade deve ser compartilhada entre médico, instituição de saúde e desenvolvedores de tecnologia apenas quando houver nexo causal demonstrado. Como enfatiza Tartuce (2022), na sociedade do risco tecnológico, a responsabilidade não pode ser afastada, apenas redistribuída conforme o grau de participação de cada agente na produção do dano.

Portanto, a IA não desfaz o paradigma da responsabilidade civil médica, e sim o redefine à luz da complexidade tecnológica atual. A lacuna legislativa brasileira sobre IA não significa ausência de responsabilidade, mas aplicação sistemática do direito vigente à nova realidade tecnológica. O desafio atual consiste em compatibilizar inovação com segurança, eficiência

com prudência e tecnologia com humanidade pois, como afirma Jonas (2006), quanto maior o poder da técnica, maior deve ser nossa responsabilidade ética para com a vida humana. Com a IA integrada à prática médica, esse dever de cuidado ganha nova dimensão. O médico que utiliza recomendações clínicas baseadas em IA passa a atuar em conjunto com os algoritmos, exigindo análise de responsabilidade mais acentuada.

A hipótese inicial se confirma: a falta de parâmetros normativos robustos amplia a vulnerabilidade de pacientes e profissionais, dificultando a responsabilização e elevando os riscos clínicos. É necessário construir um marco regulatório que assegure qualidade de dados, explicabilidade dos algoritmos, supervisão humana obrigatória e proteção integral da privacidade. O Brasil pode inspirar-se no modelo europeu do AI Act e nas diretrizes setoriais da FDA, adaptando-os à sua realidade institucional.

RISCOS E VIESES ALGORÍTMICOS EM SAÚDE

A incorporação da IA na prática clínica tem sido apresentada como uma solução capaz de ampliar a eficiência diagnóstica, reduzir custos e personalizar tratamentos médicos. Contudo, paralelamente aos benefícios, emergem riscos significativos relacionados à segurança do paciente, à justiça distributiva, à discriminação automatizada e à opacidade decisória dos algoritmos. Tais riscos não são meramente tecnológicos, possuem consequências éticas e jurídicas diretas, especialmente quando relacionados ao direito fundamental à saúde, à dignidade humana e à responsabilidade civil por danos decorrentes de decisões automatizadas.

Outro risco relevante é a opacidade algorítmica, expressão cunhada por Floridi (2018) e também conhecida como black box problem, que descreve a dificuldade ou impossibilidade de compreender o raciocínio interno de modelos complexos de machine learning e deep learning. Essa característica viola diretamente o direito à informação previsto no art. 6º, III, do Código de Defesa do Consumidor, além de comprometer o consentimento informado, que

exige compreensão do método terapêutico adotado. Em termos bioéticos, a opacidade algorítmica compromete a autonomia do paciente, pois impede que ele tome decisões com base em informação clara e compreensível (Beauchamp; Childress, 2013).

Há ainda o risco associado ao fenômeno da automação acrítica ou “automation bias”, pelo qual médicos tendem a confiar excessivamente nas recomendações do sistema, mesmo diante de inconsistências clínicas evidentes. Esse fenômeno foi identificado pela OMS, que, em relatório oficial, afirma que “a dependência cega de sistemas automatizados constitui risco ético crítico para a segurança do paciente” (WHO, 2021). A confiança exacerbada na máquina reduz a vigilância humana favorecendo erros de diagnóstico, o que, juridicamente, caracteriza negligência médica por violação do dever de prudência (Cavalieri Filho, 2019).

Existe também a possibilidade de erros de design algorítmico ou falhas de modelagem matemática, especialmente em sistemas de aprendizado profundo. Russell e Norvig (2016) descrevem essa limitação sob a forma do problema do “overfitting”, quando o algoritmo aprende o “ruído” dos dados de treino e não generaliza adequadamente para novos cenários clínicos, produzindo decisões equivocadas em ambiente real. Em medicina, isso pode significar diagnóstico incorreto, atraso no tratamento e dano grave ao paciente situação que enquadra responsabilidade civil com base no art. 927 do Código Civil.

Situações reais demonstram a gravidade desse fenômeno. Inclusive houve uma parceria entre o Google DeepMind e o Royal Free Hospital, em Londres, foi investigada por uso ilegal de dados médicos sem consentimento de mais de 1,6 milhão de pacientes (Floridi, 2018). Em 2023, pesquisadores da Universidade de Stanford alertaram que o sistema ChatGPT apresentou alucinações diagnósticas ao inventar casos clínicos em laudos médicos simulados, demonstrando risco para uso clínico sem supervisão (Stanford Medicine Report, 2023).

No Brasil, ainda não há jurisprudência específica sobre responsabilidade civil por erro médico decorrente de IA, mas o ordenamento jurídico já fornece bases normativas para enqua-

drar riscos algorítmicos. A responsabilidade civil pode ser configurada com base no dever de cautela técnica e no princípio da prevenção e precaução, amplamente aceitos no biodireito (Diniz, 2015). Desse modo, os riscos e vieses algorítmicos representam verdadeira transformação no modelo de responsabilização profissional na medicina, exigindo a formulação de protocolos éticos de uso de IA que assegurem transparência, explicabilidade, não discriminação e supervisão humana obrigatória. Como afirma Hans Jonas (2006), “quanto maior o poder tecnológico, maior deve ser a responsabilidade humana”. Na saúde, isso significa que a IA só pode ser legitimamente utilizada se estiver a serviço da dignidade da pessoa humana, fundamento do Estado Democrático de Direito (CF/88, art. 1º, III).

Assim, os riscos da IA não podem ser entendidos como meras falhas técnicas inevitáveis, mas como problemas morais e jurídicos, pois podem causar danos a pacientes, violar direitos fundamentais e comprometer a confiança na medicina. O desafio colocado ao direito é construir parâmetros normativos claros para governar algoritmos na saúde sem impedir a inovação, mas garantindo que ela permaneça sob controle humano, ético e responsável.

SÍNTESE COMPARADA (AI ACT/UE E FDA/EUA)

A ausência de regulamentação para o uso da IA na saúde no Brasil contrasta com a evolução normativa observada em outras jurisdições, especialmente na União Europeia e nos Estados Unidos. O estudo comparado permite identificar modelos regulatórios já consolidados e que podem servir de base para a construção de parâmetros jurídicos brasileiros. Tanto o AI Act da União Europeia quanto as normas do FDA dos EUA adotam a IA como tecnologia de risco e estabelecem estruturas de governança destinadas a proteger usuários contra danos técnicos, éticos e jurídicos. Esses modelos dialogam diretamente com princípios do biodireito, como precaução, segurança e promoção da dignidade humana.

O AI Act (Artificial Intelligence Act), aprovado pelo Parlamento Europeu em 2024, representa o primeiro marco regulatório abrangente do

mundo voltado especificamente para a IA. Sua abordagem normativa baseia-se no princípio da gestão de risco (risk-based approach), classificando sistemas de IA em quatro categorias: risco mínimo, risco limitado, alto risco e risco inaceitável. As aplicações da IA em saúde, sobretudo em diagnóstico, suporte clínico e prognóstico, são classificadas como IA de alto risco (AI Act, Art. 6º), exigindo rigorosos requisitos de segurança, qualidade e supervisão humana obrigatória.

Entre as exigências principais impostas pelo AI Act estão: transparência algorítmica, documentação técnica auditável, gestão de vieses, garantia de explicabilidade, rastreabilidade das decisões automatizadas e certificação prévia antes da comercialização no mercado europeu. Importante destacar que o AI Act reconhece a necessidade de preservar a autonomia do profissional de saúde, proibindo que decisões clínicas sejam tomadas exclusivamente com base em IA sem possibilidade de revisão humana (AI Act, art. 14, 2024). Isso confirma a posição europeia de que IA clínica é ferramenta de apoio e não de substituição do raciocínio médico, princípio convergente com a bioética principialista de Beauchamp e Childress (2013).

Nos Estados Unidos, o FDA adota uma abordagem diferente, focada na regulação como tecnologia médica (Software as a Medical Device, SaMD), tratada como dispositivo médico sujeito à aprovação federal. Em vez de um marco geral como o AI Act, o modelo norte-americano é pragmático e baseado na certificação técnica e no controle pós-mercado. Para IA adaptativa ou que aprende continuamente (machine learning adaptativo), o FDA exige planos de controle de risco e monitoramento constante, com possibilidade de retirada do software do mercado em caso de risco grave ou dano comprovado ao paciente (FDA, 2022).

Enquanto a União Europeia enfatiza uma abordagem preventiva e baseada em direitos fundamentais, os Estados Unidos adotam um modelo de regulação responsiva, centrado na inovação e na responsabilidade pós-fato. Assim, o FDA privilegia a análise de risco-benefício e a proteção do mercado contra produtos inseguros, alinhando-se ao pragmatismo jurídico norte-americano. Entretanto, diferente-

mente do AI Act, a regulação do FDA não apresenta diretrizes expressas sobre vieses algorítmicos ou proteção contra discriminação estrutural, o que tem suscitado críticas de autores como O'Neil (2016) e Buolamwini e Gebru (2019), que defendem a necessidade de um controle ético mais rigoroso e de mecanismos de auditoria para mitigar riscos discriminatórios nos Estados Unidos.

Em termos comparativos, é possível observar que o AI Act assume a IA em saúde como objeto de bioética normativa, unindo princípios da justiça algorítmica e direitos humanos, enquanto o FDA enxerga a IA como instrumento de tecnologia médica sujeito à comprovação empírica de eficácia e segurança. Traduzindo para o campo jurídico brasileiro, ambos os modelos trazem lições essenciais:

1. Da União Europeia: a importância da governança algorítmica, proteção aos pacientes contra riscos invisíveis e supervisão humana obrigatória;
2. Dos EUA: a necessidade de avaliação técnica e fiscalização contínua para evitar danos concretos decorrentes de falhas de software clínico.

À luz desse panorama, percebe-se que o Brasil ainda necessita de um arcabouço normativo equivalente. A LGPD (Lei nº 13.709/2018) traz diretrizes relevantes, especialmente no direito à revisão de decisões automatizadas (art. 20), mas não enfrenta a complexidade da IA em saúde em sua dimensão ética, técnica e biomédica. Da mesma forma, as resoluções do CFM sobre telemedicina e prática clínica não abordam a responsabilidade do profissional frente a sistemas algorítmicos. Urge, portanto, que o Brasil desenvolva estrutura regulatória inspirada no AI Act e no FDA, para que possa conciliar inovação com proteção jurídica do paciente. Assim, a análise comparada evidencia que a regulação internacional caminha para um modelo de responsabilidade compartilhada entre médicos, instituições de saúde e desenvolvedores de IA, com ênfase na transparência técnica e no controle de riscos. Esse cenário normativo será fundamental para orientar as propostas regulatórias brasileiras que serão discutidas no próximo capítulo deste trabalho. A regulação da IA aplicada à saúde vem sendo construída por distintos caminhos

institucionais. Na União Europeia, o AI Act criou um referencial de caráter geral que classifica sistemas de IA por níveis de risco impondo requisitos (antes da entrada no mercado). Nos Estados Unidos, a FDA adota uma abordagem setorial, regulando softwares de IA como Software as a Medical Device (SaMD), com ênfase em validação clínica, gestão do ciclo de vida e monitoramento pós-mercado. Este capítulo compara os dois modelos, mobiliza a literatura internacional e extrai lições normativas úteis ao contexto brasileiro.

A principal promessa da IA na prática médica é aumentar a precisão diagnóstica e apoiar a tomada de decisão em contextos críticos. Estudos demonstram que algoritmos podem identificar padrões em exames de imagem com velocidade superior à de radiologistas humanos. Em unidades de terapia intensiva, modelos preditivos já ajudam a estimar risco de sepse ou falência de órgãos, oferecendo informações em tempo real que auxiliam médicos sobrecarregados.

Apesar dos benefícios, desafios ainda persistem. A infraestrutura tecnológica necessária para operar sistemas de IA não está disponível de forma integrada. Hospitais públicos enfrentam limitações de equipamentos, conectividade e integração de sistemas de informação. Sem essas condições, a promessa de transformação tecnológica corre o risco de se concentrar apenas em centros privados de excelência, acentuando desigualdades no acesso à saúde.

Outro problema é a resistência cultural de parte dos profissionais. Médicos temem a perda de autonomia e a responsabilização por erros cometidos com auxílio de máquinas.

O viés algorítmico é preocupação crescente. Sistemas treinados em bases de dados não representativas podem reproduzir discriminações raciais e socioeconômicas. O resultado pode ser a exclusão do acesso a diagnósticos ou tratamentos adequados, em afronta direta ao princípio da justiça bioética. Isso reforça a necessidade de auditorias independentes e de políticas de governança de dados.

PROPOSTAS NORMATIVAS E CRITÉRIOS DE CONFORMIDADE

A ausência de regulamentação para o uso da IA na prática médica brasileira evidencia um quadro de insegurança jurídica que impacta diretamente três pilares fundamentais da ordem constitucional: a dignidade da pessoa humana (art. 1º, III, CF), o direito social à saúde (arts. 6º e 196, CF) e a segurança jurídica (art. 5º, caput, CF). Na prática, a utilização de sistemas inteligentes em ambientes clínicos vem ocorrendo sem critérios técnicos uniformes, sem supervisão normativa adequada e, sobretudo, sem definição clara de responsabilidades. Nesse cenário, a carência de diretrizes jurídicas favorece riscos éticos, clínicos e informacionais, o que torna urgente a construção de um marco regulatório nacional para a IA em saúde. Assim, este capítulo apresenta propostas normativas com base em experiências internacionais reconhecidas, como o AI Act da União Europeia e diretrizes do FDA americano, adaptadas à realidade brasileira e compatíveis com os princípios constitucionais e bioéticos que regem a prática médica.

A incorporação da IA na prática médica brasileira demanda a construção de um marco regulatório capaz de enfrentar os desafios éticos e jurídicos apresentados por essa tecnologia. Diferentemente de outras inovações em saúde, a IA interfere diretamente na tomada de decisão clínica e pode influenciar condutas terapêuticas e diagnósticas, motivo pelo qual não pode ser tratada como mero recurso auxiliar neutro.

O primeiro passo é reconhecer que a IA deve ser tratada, juridicamente, como tecnologia de risco elevado, sobretudo diante de seu potencial para causar danos significativos se utilizada de forma inadequada. Assim, deve prevalecer o princípio da precaução tecnológica, que impõe prudência e controle quando se está diante de sistemas cujos efeitos e limites ainda não são totalmente conhecidos.

A experiência internacional demonstra que é possível compatibilizar inovação tecnológica com responsabilidade jurídica. A União Europeia, por meio do AI Act, estabeleceu parâmetros rigorosos para o uso de sistemas de IA, classificando suas aplicações conforme o grau de risco e exigindo transparência, explicabili-

dade e supervisão humana obrigatória. Esse modelo europeu afirma, com clareza, que a decisão clínica não pode ser delegada exclusivamente a algoritmos, pois a responsabilidade final deve permanecer com o profissional de saúde. Já os Estados Unidos, embora adotem abordagem menos principiológica, também tratam a IA como dispositivo médico passível de certificação e auditoria pelo FDA, exigindo monitoramento contínuo do desempenho clínico dos algoritmos.

O Brasil pode e deve inspirar-se nesses modelos, sem, no entanto, simplesmente replicá-los. A construção de uma regulação própria exige o diálogo com a realidade do Sistema Único de Saúde (SUS), com a legislação já existente e com a cultura jurídica brasileira. A LGPD, por exemplo, fornece base normativa importante ao garantir ao paciente o direito de revisão de decisões automatizadas e ao impor transparência no tratamento de dados sensíveis, especialmente aqueles relacionados à saúde. No entanto, essa proteção ainda é insuficiente diante de riscos mais amplos, como viés algorítmico, opacidade técnica e possíveis erros médicos amplificados por tecnologia.

Nesse cenário, torna-se necessário consolidar diretrizes que assegurem a supervisão humana contínua no uso clínico da IA e estabeleçam deveres de auditoria técnica para os sistemas utilizados em diagnóstico e terapêutica. A transparência decisória deve ser reforçada, garantindo ao paciente o direito de saber quando a IA foi utilizada em seu caso e de compreender, de forma acessível, os fundamentos clínicos daquela decisão. Da mesma forma, é imprescindível exigir que hospitais e operadoras de saúde mantenham registros claros quando houver assistência automatizada, para que eventuais responsabilidades possam ser corretamente apuradas.

Paralelamente, é indispensável impedir qualquer forma de discriminação algorítmica que possa violar os princípios constitucionais da igualdade e da dignidade humana. A ética médica também não pode ser afastada desse debate: conforme o Código de Ética Médica, é dever do médico apenas utilizar tecnologias validadas cientificamente e cujo risco seja conhecido, não podendo transferir ao algoritmo a essência de seu julgamento profissional. Em

outras palavras, a IA deve permanecer como instrumento de apoio e jamais como substituto do raciocínio clínico. Por tudo isso, impõe-se a construção de um marco regulatório nacional para a IA em saúde, com base em responsabilidade proporcional, governança algorítmica e proteção do paciente. A tecnologia deve estar a serviço do ser humano, e não o contrário. O direito brasileiro não pode assistir passivamente à transformação da medicina pela automação, sob pena de permitir que decisões sem alma e sem responsabilidade substituam o compromisso histórico da prática médica com a vida e com a dignidade.

O Brasil discute atualmente o PL 2338/2023, que institui um marco geral para IA, e o PL 21/2020, que visa estabelecer princípios para a IA no país. No entanto, ambos não tratam da IA na saúde de forma específica, deixando sem resposta temas essenciais como responsabilidade civil médica, validação de algoritmos clínicos e proteção de dados sensíveis. Assim, impõe-se a criação de uma regulamentação setorial, como já ocorre com a FDA nos EUA e o AI Act na União Europeia, para disciplinar o uso médico da IA com foco em segurança, transparência e responsabilidade.

CONCLUSÃO

A incorporação da IA na prática médica representa um dos maiores desafios contemporâneos para o Direito, para a bioética e para a própria estrutura de proteção dos direitos fundamentais no Brasil. Longe de significar apenas um avanço tecnológico, a IA transforma a natureza do ato médico inaugurando uma nova relação entre técnica, responsabilidade e cuidado humano. Como demonstrado neste trabalho, a tecnologia não é neutra: carrega valores, prioridades e escolhas de design que podem impactar diretamente a vida, a integridade física e a dignidade dos pacientes. A análise desenvolvida ao longo deste trabalho permitiu concluir que a ausência de uma normatização específica para o uso da IA na prática médica brasileira gera riscos concretos e não pode mais ser ignorada pelo Direito. Embora a IA represente um avanço relevante para a medicina contemporânea e ofereça benefícios inegáveis, como maior precisão diagnós-

tica e agilidade na análise de dados clínicos complexos, sua adoção desregulada compromete valores constitucionais essenciais, entre eles a dignidade da pessoa humana, a proteção dos dados pessoais sensíveis e a própria segurança jurídica da atividade médica.

Verificou-se que o ordenamento jurídico brasileiro ainda não dispõe de uma regulação adequada para o uso da IA em saúde, limitando-se a aplicar, de forma adaptada, normas gerais do Código Civil, do Código de Defesa do Consumidor, do Código de Ética Médica e da LGPD. Embora esses dispositivos ofereçam respostas parciais, não são suficientes para enfrentar questões inéditas trazidas pelos sistemas algorítmicos, tais como opacidade decisória, viés matemático, responsabilidade compartilhada, tomada de decisão automatizada e autonomia informacional do paciente. Fica evidente que o uso crescente de sistemas de apoio à decisão clínica, ainda que sob a justificativa de modernização e eficiência, ameaça a autonomia do paciente quando não há transparência quanto à participação de algoritmos no processo diagnóstico ou terapêutico. Além disso, a opacidade técnica de muitos sistemas algorítmicos esvazia o dever de informação e prejudica o consentimento esclarecido, desumanizando a relação médico-paciente ao substituir a decisão clínica fundamentada por resultados automatizados de justificativa limitada ou inexistente.

Ao mesmo tempo, a pesquisa evidenciou que a ausência de diretrizes normativas também afeta os profissionais de saúde, que passam a atuar em um cenário de insegurança jurídica, sem clareza quanto à extensão de sua responsabilidade civil em casos de falha ou dano decorrente do uso de IA. A falta de regulamentação abre margem para conflitos de responsabilidade entre médicos, instituições de saúde e empresas desenvolvedoras de software, dificultando a reparação do paciente lesado e fragilizando a confiança na prática clínica.

Dessa forma, a IA pode e deve integrar a prática médica contemporânea, desde que utilizada de maneira ética, transparente e juridicamente controlada. O Direito não deve impedir o avanço tecnológico, mas orientá-lo para que esteja sempre a serviço da vida humana e da justiça. Regular a IA em saúde, portanto, não significa frear o progresso, mas assegurar que

ele caminhe de forma compatível com os valores fundamentais que estruturam o Estado Democrático de Direito.

Ficou demonstrado também que a responsabilidade civil médica permanece válida e necessária mesmo diante da IA, que deve ser compreendida juridicamente como instrumento tecnológico de apoio, e não como substituto do julgamento clínico. O médico continua responsável pelas escolhas terapêuticas, devendo agir com prudência técnica reforçada, especialmente diante de riscos ou limitações do sistema utilizado. A confiança cega em algoritmos, como alertado, constitui negligência profissional e afronta ao dever de cuidado.

Ao mesmo tempo, a IA traz riscos que extrapolam a esfera do erro médico tradicional. Estudos recentes comprovaram a existência de vieses algorítmicos capazes de gerar discriminação, desigualdade de acesso e decisões injustas em saúde, reforçando desigualdades sociais pré-existentes. Assim, a discussão jurídica deve assumir uma dimensão ética: não basta responsabilizar danos, é necessário evitá-los, adotando medidas de prevenção e governança tecnológica. Nesse ponto, a convergência entre Direito e bioética revelou-se indispensável, especialmente com base nos princípios da autonomia, beneficência, não maleficência e justiça.

Diante disso, defendeu-se a necessidade de um marco regulatório para a IA em saúde no Brasil, inspirado nas experiências internacionais como o AI Act da União Europeia e as diretrizes do FDA americano. Esse marco deve incluir supervisão humana obrigatória, mecanismos de auditoria algorítmica, responsabilidade solidária entre médico, hospital e desenvolvedor, controle de vieses e garantia de transparência decisória. Ao paciente deve ser assegurado o direito de informação, o direito de contestação e o direito de revisão humana de decisões automatizadas.

Conclui-se que a IA pode, sim, fortalecer a medicina, desde que submetida a limites ético-jurídicos que impeçam sua instrumentalização contra a vida humana. A tecnologia deve permanecer subordinada ao Direito e à ética médica, e nunca acima deles. Se o século XXI exige inovação, exige também responsabilidade. A ciência avança, mas a dignidade humana permanece como fundamento da República

e limite inegociável do progresso. Assim, a regulação da IA na saúde não é uma escolha: é uma exigência constitucional e civilizatória.

Portanto, regulamentar a IA na medicina não é uma opção futurista, é uma exigência constitucional imediata. O Brasil deve estabelecer um marco regulatório específico para IA em saúde, inspirado em boas práticas internacionais, com equilíbrio entre inovação, ética e proteção jurídica. Somente assim será possível garantir que a tecnologia avance a serviço da vida e da dignidade humana, e não contra elas.

REFERÊNCIAS

BEAUCHAMP, Tom L.; CHILDRESS, James F. **Principles of Biomedical Ethics**. 7th. New York: Oxford University Press, 2013.

BRASIL. **Constituição da República Federativa do Brasil de 1988**. Brasília, DF: Senado Federal, 1988.

BRASIL. **Lei n. 13.709, de 14 de agosto de 2018**. Lei Geral de Proteção de Dados Pessoais (LGPD). Diário Oficial da União, Brasília, DF, 15 ago. 2018.

BUOLAMWINI, Joy; GEBRU, Timnit. **Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification**. Cambridge, MA: MIT Media Lab, 2019. Disponível em: <http://gendershades.org>. Acesso em: 22 out. 2025.

CAVALIERI FILHO, Sérgio. **Programa de responsabilidade civil**. 14. ed. São Paulo: Atlas, 2019.

CONSELHO FEDERAL DE MEDICINA (Brasil). **Resolução CFM n. 2.217/2018**. Código de Ética Médica. Brasília, DF: CFM, 2018.

CONSELHO FEDERAL DE MEDICINA (Brasil). **Resolução CFM n. 2.314/2022**. Define e regulamenta a telemedicina, como forma de serviços médicos mediados por tecnologias de comunicação. Brasília, DF: CFM, 2022. Disponível em: <https://sistemas.cfm.org.br/normas/arquivos/>

resolucoes/BR/2022/2314_2022.pdf. Acesso em: 12 set. 2025.

DINIZ, Maria Helena. **O estado atual do bio-direito**. 9. ed. São Paulo: Saraiva, 2015.

DONEDA, Danilo. **Da privacidade à proteção de dados pessoais**: elementos da formação da autoridade e da dignidade da pessoa na sociedade da informação. 2006. Tese (Doutorado em Direito Civil) – Universidade do Estado do Rio de Janeiro, Rio de Janeiro, 2006.

DONEDA, Danilo. **Da privacidade à proteção de dados pessoais**: elementos da formação da autoridade contemporânea. 2. ed. Rio de Janeiro: Forense, 2011.

FLORIDI, Luciano. Soft Ethics and the Governance of the Digital. **Philosophy & Technology**, v. 31, p. 1-8, mar. 2018.

JONAS, Hans. **O princípio responsabilidade**: ensaio de uma ética para a civilização tecnológica. Rio de Janeiro: Contraponto, 2006.

MITTELSTADT, B. D.; ALLO, P.; TADDEO, M.; WACHTER, S.; FLORIDI, L. The ethics of algorithms: Mapping the debate. **Big Data & Society**, v. 3, n. 2, p. 1–21, 2016. Disponível em: <https://doi.org/10.1177/2053951716679679>. Acesso em: 22 out. 2025.

NOBLE, Safiya Umoja. **Algorithms of Oppression**: How Search Engines Reinforce Racism. New York: NYU Press, 2018.

OBERMEYER, Z.; POWERS, B.; VOGELI, C.; MULLAINATHAN, S. Dissecting racial bias in an algorithm used to manage the health of populations. **Science**, v. 366, n. 6464, p. 447–453, 2019. DOI: 10.1126/science.aax2342. Acesso em: 22 out. 2025.

O'NEIL, Cathy. **Weapons of Math Destruction**: How Big Data Increases Inequality and Threatens Democracy. New York: Crown Publishing, 2016.

POTTER, Van Rensselaer. **Bioethics: Bridge to the Future**. Englewood Cliffs, NJ: Prentice-Hall, 1971.

RUSSELL, Stuart; NORVIG, Peter. **Artificial Intelligence: A Modern Approach**. 3rd ed. Upper Saddle River, NJ: Prentice Hall, 2016.

STANFORD MEDICINE. **Evaluating Large Language Models in Clinical Decision Support**: Early Findings from ChatGPT in Medicine. Stanford, CA: Stanford University School of Medicine, 2023. Disponível em: <https://med.stanford.edu/>. Acesso em: 22 out. 2025.

TARTUCE, Flávio. **Manual de direito civil: responsabilidade civil**. 12. ed. São Paulo: Método, 2022.

TOPOL, Eric. **Deep Medicine**: How Artificial Intelligence Can Make Healthcare Human Again. New York: Basic Books, 2019.

UNIÃO EUROPEIA. Regulamento (UE) 2016/679 do Parlamento Europeu e do Conselho, de 27 de abril de 2016. Relativo à proteção das pessoas singulares no que diz respeito ao tratamento de dados pessoais e à livre circulação desses dados (Regulamento Geral sobre a Proteção de Dados – GDPR). **Jornal Oficial da União Europeia**, L 119, 4 May 2016.

UNIÃO EUROPEIA. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 on Artificial Intelligence (Artificial Intelligence Act). **Official Journal of the European Union**, L 202, 12 jul. 2024. Disponível em: <https://eur-lex.europa.eu/eli/reg/2024/1689>. Acesso em: 22 out. 2025.

UNITED STATES FOOD AND DRUG ADMINISTRATION (FDA). **Artificial Intelligence/ Machine Learning (AI/ML)-Based Software as a Medical Device (SaMD) Action Plan**. Silver Spring, MD: U.S. FDA, 2022. Disponível em: <https://www.fda.gov/medical-devices/software-medical-device-samd/artificial-intelligence-and-machine-learning-software-medical-device>. Acesso em: 22 out. 2025.



WORLD HEALTH ORGANIZATION. **Ethics and governance of artificial intelligence for health**: WHO guidance. Geneva: World Health Organization, 2021. Disponível em: <https://www.who.int/publications/i/item/9789240029200>. Acesso em: 22 out. 2025.