

Regulamentação da Inteligência Artificial na União Europeia: estrutura ética, classificação de riscos e possíveis reflexos na medicina¹

Regulación de la Inteligencia Artificial en la Unión Europea: estructura ética, clasificación de riesgos y posibles consecuencias en medicina

Regulation of Artificial Intelligence in the European Union: ethical framework, risk classification and possible consequences in medicine

Alexandre Rocha Almeida de Moraes²

Doutor, PPG em Direito da Saúde, Universidade Santa Cecília, Santos, SP, Brasil

ORCID: <https://orcid.org/0000-0002-8374-5694>

Andressa Felix Lisboa³

Mestranda, PPG em Direito da Saúde, Universidade Santa Cecília, Santos, SP, Brasil

ORCID: <https://orcid.org/0009-0001-0585-4828>

RESUMO: Contextualização: A regulamentação da inteligência artificial pela União Europeia busca estabelecer um quadro claro para a classificação de riscos associados a esses sistemas, abrangendo categorias de risco mínimo, limitado, alto e inaceitável. **Problema:** O problema central reside na necessidade de garantir que esses sistemas não comprometam a segurança e os direitos fundamentais das pessoas. Por outro lado, o excesso de regulação, além de inviabilizar para parte do Ocidente o desenvolvimento tecnológico em áreas como a medicina, pode se tornar ineficaz frente à competitividade das grandes potências mundiais. **Objetivos:** Este estudo visa descrever a metodologia de classificação de riscos e os princípios éticos que orientam a implementação de um modelo de inteligência artificial confiável. **Métodos:** Utilizou-se uma revisão narrativa documental do regulamento europeu, com análise qualitativa para identificar e categorizar os requisitos específicos de cada nível de risco da Inteligência Artificial, confrontando as vantagens e desvantagens da revolução tecnológica à medicina. **Resultados:** Os sistemas de Inteligência Artificial classificados como de risco mínimo, como filtros de spam, são considerados seguros e não exigem supervisão rigorosa. Já os sistemas de alto risco, como Inteligência Artificial para diagnóstico médico, requerem conformidade rigorosa com padrões éticos e de segurança, devido ao potencial de impacto sobre a vida humana. De qualquer sorte, frente aos potenciais avanços da tecnológica para a medicina e para a saúde das pessoas, coloca-se a questão ética de se aferir quanto a regulação jurídica pode atrasar novas descobertas para a saúde individual e coletiva. **Considerações Finais:** As considerações finais indicam que a proposta regulamentar contribui significativamente para a construção de um modelo de segurança, mas enquanto essa regulação não for global, pode ser inócua ou até mesmo obstar o desenvolvimento científico de áreas como a medicina.

Palavras-chave: Inteligência Artificial, Medição de Risco, Direito da Saúde

¹ Esse trabalho foi apresentado originalmente no VI Congresso Internacional de Direito da Saúde e XII Congreso Iberoamericano de Derecho Sanitario, realizado em 16, 17, 18 e 19 de outubro de 2024 na Universidade Santa Cecília (UNISANTA). Em função da recomendação de publicação da Comissão Científica do Congresso, fez-se a presente versão. Menção honrosa de melhor trabalho apresentado no Congresso, para o Grupo de Trabalho “Novas Tecnologias e o Direito da Saúde”.

² Promotor de Justiça do Ministério Público do Estado de São Paulo. Mestre e Doutor em Direito pela Pontifícia Universidade Católica de São Paulo. É professor da Graduação e Pós-Graduação da PUC/SP e da UNISANTA. Professor do Mestrado em Direito da Saúde (UNISANTA), professor de diversos cursos de pós-graduação, dentre os quais a Escola Superior do Ministério Público de São Paulo e da Escola Paulista da Magistratura. Autor de obras jurídicas, dentre as quais, Direito Penal: Parte Geral (Fórum), Direito Penal do Inimigo e Direito Penal Racional (Editora Juruá) e Criminologia (Juspodvum).

³ Advogada, especialista em seguridade social. Mestranda em Direito da Saúde pela UNISANTA. Bolsista CAPES. Atua como professora em cursos de extensão sobre previdência social pela ESA/SP. Professora convidada em direito do trabalho para cursos de extensão EAD - UNISANTA. Responsável pela área previdenciária do escritório Lamy & Oliveira sociedade de advogados. Membro efetiva da comissão de estágio e exame de ordem OAB/SP.

DOI: 10.5281/zenodo.14174188

RESUMEN: Contextualización: La regulación de la inteligencia artificial por la Unión Europea busca establecer un marco claro para la clasificación de riesgos asociados a estos sistemas, abarcando categorías de riesgo mínimo, limitado, alto e inaceptable. **Problema:** El problema central reside en la necesidad de garantizar que estos sistemas no comprometan la seguridad y los derechos fundamentales de las personas. Por otro lado, el exceso de regulación, además de dificultar el desarrollo tecnológico en áreas como la medicina para parte del Occidente, puede volverse ineficaz frente a la competitividad de las grandes potencias mundiales. **Objetivos:** Este estudio tiene como objetivo describir la metodología de clasificación de riesgos y los principios éticos que orientan la implementación de un modelo de inteligencia artificial confiable. **Métodos:** Se utilizó una revisión narrativa documental de la regulación europea, con análisis cualitativo para identificar y categorizar los requisitos específicos para cada nivel de riesgo de Inteligencia Artificial, comparando las ventajas y desventajas de la revolución tecnológica en medicina. **Resultados:** Los sistemas de Inteligencia Artificial clasificados como de mínimo riesgo, como los filtros de spam, se consideran seguros y no requieren una estrecha supervisión. Los sistemas de alto riesgo, como la Inteligencia Artificial para el diagnóstico médico, requieren un estricto cumplimiento de estándares éticos y de seguridad, debido al potencial impacto en la vida humana. En cualquier caso, ante los potenciales avances tecnológicos en la medicina y la salud de las personas, surge la cuestión ética de valorar hasta qué punto la regulación legal puede retrasar nuevos descubrimientos para la salud individual y colectiva. **Consideraciones finales:** Las consideraciones finales indican que la propuesta regulatoria contribuye significativamente a la construcción de un modelo de seguridad, pero, mientras esta regulación no sea global, puede resultar inofensiva o incluso obstaculizar el desarrollo científico en áreas como la medicina.

Palabras clave: Inteligencia Artificial, Medición de Riesgo, Derecho Sanitario.

ABSTRACT: Contextualization: The European Union's regulation of artificial intelligence seeks to establish a clear framework for classifying risks associated with these systems, covering categories of minimal, limited, high, and unacceptable risk. **Problem:** The central issue lies in ensuring that these systems do not compromise the safety and fundamental rights of individuals. On the other hand, excessive regulation, in addition to hindering technological development in fields such as medicine in some parts of the West, may become ineffective in the face of the competitiveness of major world powers. **Objectives:** This study aims to describe the risk classification methodology and the ethical principles guiding the implementation of a reliable artificial intelligence model. **Methods:** A documentary narrative review of the European regulation was used, with qualitative analysis to identify and categorize the specific requirements for each level of Artificial Intelligence risk, comparing the advantages and disadvantages of the technological revolution in medicine. **Results:** Artificial Intelligence systems classified as minimal risk, such as spam filters, are considered safe and do not require close supervision. High-risk systems, such as Artificial Intelligence for medical diagnosis, require strict compliance with ethical and safety standards, due to the potential impact on human life. In any case, given the potential technological advances in medicine and people's health, the ethical question arises of assessing how much legal regulation can delay new discoveries for individual and collective health. **Final considerations:** The final considerations indicate that the regulatory proposal significantly contributes to building a security model, but as long as this regulation is not global, it may be ineffective or even hinder scientific development in areas such as medicine.

Keywords: Artificial Intelligence, Risk Assessment, Health Law.

Introdução

A transição da manufatura para as máquinas, assim como a invenção do computador, a revolução tecnológica e dos meios de comunicação representam verdadeiros marcos na história.

A transição da manufatura à indústria mecânica, a introdução de máquinas fabris

DOI: 10.5281/zenodo.14174188

e o crescimento do rendimento do trabalho aumentaram a produção global.

A invenção das máquinas desencadeou a produção em série, acelerou a circulação das mercadorias, gerando como subprodutos a divisão do trabalho, a urbanização, os avanços das ciências e, em especial, da medicina preventiva e sanitária, além do controle de epidemias que aumentaram a expectativa de vida e favoreceram o crescimento demográfico (Moraes, p. 37).

Transformaram-se as demandas e os problemas sociais. Os riscos de morte prematura deram lugar a uma população que passou a gerar novos conflitos, inclusive de criminalidade. A mesma máquina criada para o progresso serviu, por exemplo, para matar em acidentes automobilísticos. A industrialização e os meios de transporte alçaram novos bens para o Direito Penal. Os riscos advindos das máquinas e dos meios de transporte, por certo, geraram novas demandas à vida e à integridade física das pessoas que, até então, não existiam. (Moraes, p. 38). Da mesma maneira, a revolução tecnológica vem desencadeando alterações políticas, econômicas e sociais profundas: novos atores sociais, novos costumes e uma nova sociedade. A próxima revolução já se iniciou: a inteligência artificial generativa e a robótica inteligente.

O objetivo central deste estudo é investigar como o Regulamento Europeu sobre Inteligência Artificial aborda a questão do viés e da discriminação em sistemas de IA e avaliar a eficácia dessas diretrizes na mitigação desses problemas.

Ao mesmo tempo, imaginando o quanto a saúde e, em especial, a medicina tende a evoluir nos próximos anos com o uso da IA, o presente estudo coloca, de maneira zetética, um questionamento: a proposta de regulamentar contribui significativamente para a construção de um modelo de segurança, mas enquanto essa regulação não for global, não será inócua diante da disponibilidade de serviços e produtos estrangeiros não regulados e, em especial, não poderia implicar em um retardo no desenvolvimento científico de áreas como a medicina para o Ocidente?

A presente pesquisa valeu-se da revisão narrativa (Lamy, 2020, pp. 338-339) sobre a regulação da Inteligência Artificial (IA) pelo Parlamento Europeu, com especial atenção aos princípios éticos estabelecidos pela União Europeia. A pesquisa incluiu uma análise bibliográfica de textos acadêmicos sobre princípios éticos em direito e tecnologia, bem como uma investigação documental de instrumentos jurídicos internacionais. A isso, agregou-se algumas análises críticas dos autores do presente texto.

Dessa forma, este estudo contribui para a compreensão das políticas de IA na União Europeia, fornecendo *insights* valiosos sobre suas forças, fraquezas e oportunidades para melhorias futuras.

1 A Inteligência Artificial

A origem da inteligência artificial generativa (IA generativa) está profundamente enraizada no desenvolvimento da inteligência artificial e das redes neurais. A trajetória pode ser traçada através de vários marcos importantes na história da IA e da computação.

Em 1950, Alan Turing propôs o Teste de Turing como uma forma de avaliar a inteligência de uma máquina, enquanto em 1958, Frank Rosenblatt desenvolveu o

perceptron, um dos primeiros modelos de redes neurais, que serviu como base para o desenvolvimento futuro de redes neurais mais complexas (Rosenblatt, 1958, p. 534).

Já nos anos 1980 e 1990, houve uma espécie de redescoberta dos algoritmos, permitindo o treinamento eficaz de redes neurais de multicamadas, o que foi um avanço crucial para o desenvolvimento da IA generativa. Durante este período, métodos inspirados na evolução biológica começaram a ser aplicados na geração de soluções inovadoras, incluindo a criação de novas formas de conteúdo (Rumelhart, Hinton & Williams, 1986).

O avanço do denominado deep learning (aprendizado profundo) a partir dos anos 2000, impulsionado pelo aumento do poder computacional e pela disponibilidade de grandes volumes de dados (big data), permitiu a criação de modelos de redes neurais ainda mais complexos e eficazes, já aptos ao reconhecimento de imagens e processamento de linguagem natural.

Arquiteturas específicas de redes neurais começaram a ser amplamente utilizadas para tarefas como reconhecimento de imagens e processamento de linguagem natural, respectivamente. Estes modelos de linguagem, treinados em vastas quantidades de dados, demonstraram uma capacidade sem precedentes de gerar texto coerente e criativo, popularizando ainda mais a IA generativa.

A cada mês a inteligência artificial generativa continua a evoluir, com avanços contínuos em algoritmos, arquiteturas de rede e aplicações práticas. O futuro iminente já promete ainda mais integração da IA generativa em várias indústrias, desde a criação de conteúdo até a automação criativa e o design assistido por IA.

A evolução dessa tecnologia já trouxe e ainda trará, por certo, uma série de benefícios significativos, transformando setores como saúde, educação, finanças e segurança. No entanto, essa mesma evolução também levanta questões éticas cruciais, particularmente no que diz respeito ao uso da IA para a classificação de pessoas e grupos.

Entre os desafios éticos mais prementes está a questão do viés ideológico, religioso ou de qualquer natureza alheio à ciência, assim como da discriminação, que podem ser involuntariamente perpetuados pelos sistemas de IA, quando o treinamento é insuficiente ou dolosamente dirigido nesse sentido.

A União Europeia, reconhecendo esses desafios, implementou o Regulamento Europeu sobre Inteligência Artificial (EU AI Act, 2024) com o objetivo de estabelecer diretrizes que assegurem o uso ético e seguro dessas tecnologias.

Os princípios éticos fundamentais delineados pelo Regulamento Europeu sobre Inteligência Artificial incluem a necessidade de transparência, responsabilidade, privacidade, e, crucialmente, a equidade e a não discriminação. Estudos recentes demonstram, que sistemas de IA treinados com dados enviesados podem perpetuar e até exacerbar discriminações existentes, resultando em impactos negativos substanciais para indivíduos e grupos vulneráveis.

1.1 O Regulamento Inteligência Artificial da União Europeia

O parlamento Europeu definiu a inteligência artificial (IA) como sendo o uso de tecnologias digitais para desenvolver sistemas que podem realizar funções normalmente atribuídas à inteligência humana (Parlamento Europeu, 2020).

Em outubro de 2020, o Conselho Europeu colocou como pauta a transição digital e destacou a importância da inteligência artificial (IA) para o futuro da Europa (Comissão Europeia, 2020).

O Conselho Europeu, aproveitou a pauta, e fez um apelo à Comissão Europeia para propor formas de aumentar os investimentos públicos e privados na pesquisa, inovação e implantação da IA, para assegurar uma melhor coordenação entre centros de investigação europeus e fornecer uma definição clara e objetiva dos sistemas de IA e classificando-os de acordo com o seu nível de risco para a sociedade (Conselho da União Europeia, 2020, p.6).

Em abril de 2021 em resposta ao contexto político delineado pelo Conselho Europeu, a Comissão Europeia apresentou então uma proposta para a formação do Regulamento Europeu de IA (Comissão Europeia, 2021).

A proposta buscava estabelecer objetivos específicos para o Regulamento Europeu de IA buscando responder às preocupações sobre segurança, direitos fundamentais e inovação tecnológica. A proposta visava, em primeiro lugar, assegurar que os sistemas de IA sejam seguros e alinhados com a legislação europeia, de modo a proteger direitos fundamentais e os valores da União. Além disso, garantir segurança jurídica foi considerado essencial para incentivar investimentos, permitindo um ambiente favorável à inovação. Em relação à governança, a Comissão propôs fortalecer a aplicação das leis para assegurar o cumprimento dos requisitos de segurança pelos sistemas de IA, uma necessidade crescente para acompanhar o rápido desenvolvimento tecnológico. Por fim, ao facilitar o desenvolvimento de um mercado único para aplicações de IA seguras e confiáveis, a proposta visava evitar a fragmentação do mercado europeu, promovendo um ecossistema uniforme e sustentável para a IA na União Europeia:

Garantir que os sistemas de IA colocados no mercado da União e utilizados sejam seguros e respeitem a legislação em vigor em matéria de direitos fundamentais e valores da União; **Garantir** a segurança jurídica para facilitar os investimentos e a inovação no domínio da IA; **Melhorar** a governação e a aplicação efetiva da legislação em vigor em matéria de direitos fundamentais e dos requisitos de segurança aplicáveis aos sistemas de IA; **Facilitar** o desenvolvimento de um mercado único para as aplicações de IA legítimas, seguras e de confiança e evitar a fragmentação do mercado. (Comissão Europeia, 2021) (grifo nosso).

Estes objetivos, buscavam estabelecer um quadro jurídico equilibrado e flexível, focado em uma abordagem regulamentar baseada no risco que não limita indevidamente a evolução tecnológica nem aumenta desproporcionalmente os custos de mercado.

O regulamento proposto elencou regras harmonizadas para a IA na União Europeia, incluindo restrições a práticas prejudiciais e obrigações específicas para sistemas como o de identificação biométrica à distância. Com a proposta de uma metodologia para identificar sistemas de "risco elevado", propondo requisitos obrigatórios e avaliações de conformidade antes da comercialização (Comissão Europeia, 2021).

A proposta abordou também uma governança a nível dos Estados-Membros em cooperação com o Comitê Europeu para a Inteligência Artificial, além de medidas para apoiar a inovação e reduzir encargos regulatórios para pequenas e médias empresas (PME) e *startups*. Mas, foi somente no ano de 2022 que a proposta enfrentou um intenso processo de revisão e discussão. Na ocasião, o regulamento foi submetido a um exame minucioso pelo Parlamento Europeu e pelo Conselho da União Europeia.

Esse processo incluiu consultas públicas e análises de impacto que permitiram o envolvimento de diversas partes interessadas, incluindo especialistas em tecnologia, legisladores e representantes da sociedade civil (Conselho da União Europeia, 2022, p. 1-3). O objetivo principal durante essa fase foi ajustar o regulamento para refletir melhor as preocupações emergentes e garantir que os requisitos fossem práticos e eficazes.

Em 09 de dezembro de 2023, após intensas negociações, o Conselho e o Parlamento Europeu chegaram a um acordo provisório sobre o Regulamento de Inteligência Artificial.

Mas, somente em 21 de maio de 2024, o Conselho da União Europeia aprovou formalmente o Regulamento de Inteligência Artificial, estabelecendo um marco global pioneiro na regulamentação de IA. O regulamento cria um conjunto de regras que visa garantir a segurança e a conformidade dos sistemas de IA com os princípios éticos e os direitos fundamentais da União Europeia. Com isso, buscou-se não só assegurar que os sistemas de IA respeitem esses valores, mas também promover um mercado interno mais integrado e eficiente. Outro objetivo central é apoiar a inovação de forma segura e responsável, equilibrando a proteção dos cidadãos e a criação de um ambiente de desenvolvimento tecnológico sustentável. Nesta perspectiva, o documento oficial prevê:

A finalidade do presente regulamento é melhorar o funcionamento do mercado interno mediante o estabelecimento de um quadro jurídico uniforme, em particular para o desenvolvimento, a colocação no mercado, a colocação em serviço e a utilização de sistemas de inteligência artificial (sistemas de IA) na União, em conformidade com os valores da União, a fim de promover a adoção de uma inteligência artificial (IA) centrada no ser humano e de confiança, assegurando simultaneamente um elevado nível de proteção da saúde, da segurança, dos direitos fundamentais consagrados na Carta dos Direitos Fundamentais da União Europeia ("Carta"), nomeadamente a democracia, o Estado de direito e a proteção do ambiente, contra os efeitos nocivos dos sistemas de IA na União, e de apoiar a inovação. O presente regulamento assegura a livre circulação transfronteiriça de produtos e serviços baseados em IA evitando assim que os Estados-Membros imponham restrições ao desenvolvimento, à comercialização e à utilização dos sistemas de IA, salvo se explicitamente autorizado pelo presente regulamento. (Conselho da União Europeia, 2024, p. 4).

Observa-se que o regulamento tem o intuito de equilibrar o progresso tecnológico com a proteção dos direitos dos cidadãos, estabelecendo princípios fundamentais como a segurança, a transparência e a responsabilidade. Princípios estes, que foram extraídos de diversas fontes e diretrizes pré-existentes, refletindo um esforço para garantir que o desenvolvimento e uso de IA sejam seguros, transparentes e respeitem os direitos fundamentais.

1.2 Princípios e requisitos para uma Inteligência Artificial de confiança

A Carta dos Direitos Fundamentais da União Europeia foi uma base importante para a proteção dos direitos e das liberdades fundamentais, assegurando que a IA não viole direitos humanos. Além disso, a Comissão Europeia já havia publicado em abril de 2019 diretrizes sobre a IA confiável, identificando princípios e requisitos para uma IA de confiança que foram observados pelo regulamento.

O documento à época foi denominado como "Orientações Éticas para uma IA de Confiança", e detalha quatro princípios éticos e sete requisitos essenciais para garantir a confiança na inteligência artificial (IA). Os quatro princípios éticos chave são respeito pela autonomia humana, prevenção de danos, equidade e explicabilidade.

O respeito pela autonomia humana enfatiza a importância de permitir que os seres humanos mantenham controle sobre as decisões influenciadas ou tomadas pelos sistemas de IA reforçando a capacidade dos indivíduos de fazer escolhas informadas e livres, protegendo-os de manipulação e coerção.

A prevenção de danos, por sua vez, busca centrar-se na proteção dos seres humanos contra riscos e danos causados por sistemas de IA assegurando a segurança e o bem-estar dos indivíduos e da sociedade. A equidade garante que os sistemas de IA sejam desenvolvidos e operados de maneira justa e imparcial, prevenindo discriminação e promovendo justiça.

A explicabilidade refere-se à capacidade de compreender e explicar como um sistema de IA toma decisões, assegurando transparência e permitindo que os usuários confiem nas decisões automatizadas e possam contestá-las se necessário (Orientações Éticas para uma IA de Confiança, 2019, p. 48-50).

Além dos princípios, os sete requisitos elencados pela orientação ética e utilizados no regulamento para uma IA de confiança são ação e supervisão humana, robustez e segurança técnica, privacidade e governança de dados, transparência, diversidade, não discriminação e equidade, bem-estar social e ambiental, e prestação de contas.

A ação e supervisão humana garantem que os sistemas de IA incluam mecanismos de supervisão humana para alinhar suas operações com valores e princípios humanos. A robustez e segurança técnica asseguram que os sistemas de IA sejam tecnicamente robustos e seguros, evitando falhas e ataques e garantindo funcionalidade em condições adversas. A privacidade e governança de dados protegem a privacidade e garantem a gestão adequada dos dados utilizados e gerados pelos sistemas de IA. A transparência é necessária para garantir que os sistemas de IA sejam compreensíveis e que seus processos e resultados sejam claros para os usuários. A diversidade, não discriminação e equidade promovem a inclusão e a equidade, evitando vieses e discriminação.

O bem-estar social e ambiental assegura que a IA contribua para o bem-estar social e ambiental, minimizando impactos negativos e promovendo benefícios para a sociedade. A prestação de contas envolve mecanismos claros para assegurar que os sistemas de IA sejam responsáveis por suas ações e impactos (Orientações Éticas para uma IA de Confiança, 2019, p. 52-60).

Os sete requisitos são interligados e devem ser considerados com igual importância, pois se apoiam mutuamente e garantem a confiança na IA. Eles devem ser aplicados e avaliados ao longo de todo o ciclo de vida de um sistema de IA, desde o desenvolvimento até a implementação e uso contínuo. Essa abordagem holística e

integrada assegura que os sistemas de IA operem de maneira confiável, segura e ética, promovendo a confiança dos usuários e beneficiando a sociedade como um todo (Orientações Éticas para uma IA de Confiança, 2019, p. 61).

Observa-se, portanto, que o Regulamento da Comissão Europeia, avança ao utilizar os princípios e requisitos como fundamento para utilização de uma IA de confiança, e ainda classificar os riscos associados à inteligência artificial, prevendo regras específicas para cada risco elencado.

1.3 Classificação de riscos

Com intuito expresso de garantir que os sistemas de IA de alto risco sejam sujeitos a requisitos obrigatórios rigorosos para assegurar a sua segurança e conformidade com os direitos fundamentais (Regulamento, 2024, p.6).

O regulamento destaca a importância de estabelecer uma metodologia para a classificação de riscos associados à inteligência artificial (IA). A classificação de riscos é essencial para identificar e mitigar os possíveis perigos que os sistemas de IA podem representar para a sociedade, a economia e a segurança. A metodologia sugerida envolve avaliar os riscos com base em diversos fatores, incluindo o ciclo de vida do modelo de IA, suas condições de uso, confiabilidade, equidade, segurança e o nível de autonomia (Regulamento, p. 62).

Os riscos sistêmicos são de particular preocupação, especialmente quando as capacidades de um modelo de IA de finalidade geral são de alto impacto, ou seja, quando essas capacidades excedem as de modelos mais avançados.

Esses riscos podem ser influenciados por fatores como a possibilidade de uso indevido intencional, falhas não intencionais de controle, riscos químicos, biológicos, radiológicos e nucleares, capacidades cibernéticas ofensivas, e o potencial de prejudicar infraestruturas críticas. Além disso, preocupações adicionais incluem a possibilidade de a IA gerar vieses prejudiciais e discriminação, facilitando a desinformação ou ameaçando a privacidade e os direitos humanos, podendo levar a uma reação em cadeia com efeitos negativos significativos (Regulamento, 2024, p. 64-66).

Os sistemas de IA destinados à administração da justiça e aos processos democráticos são considerados de alto risco devido ao seu impacto potencialmente significativo na democracia, no estado de direito e nas liberdades individuais. Portanto, é apropriado classificá-los como de alto risco para mitigar os potenciais vieses, erros e falta de transparência que podem surgir nesses contextos.

A decisão final em processos judiciais, no entanto, deve permanecer uma atividade humana, sem ser totalmente substituída pela IA. Essa classificação também se aplica a sistemas de IA utilizados para influenciar resultados eleitorais ou comportamentos eleitorais, a fim de proteger a integridade dos processos democráticos (Regulamento, 2024, p. 68-69).

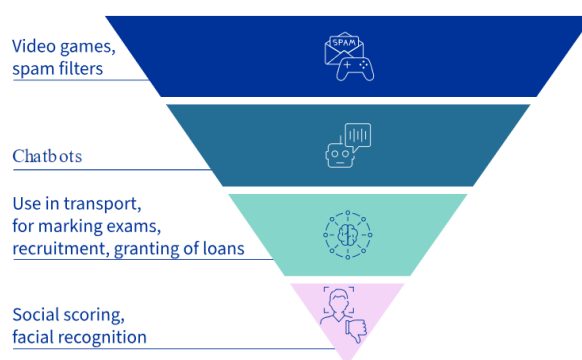
A metodologia de classificação de riscos proposta no documento é essencial para assegurar que os modelos de IA de alta capacidade e impacto sejam rigorosamente avaliados e monitorados, minimizando os potenciais perigos e maximizando os benefícios desses sistemas para a sociedade. Para isso, são necessários mecanismos robustos de

avaliação de conformidade e supervisão contínua, garantindo que os sistemas de IA operem dentro dos padrões éticos e de segurança estabelecidos (Regulamento, 2024, p. 70-72).

Em suma, nota-se que o regulamento sublinha que a classificação dos riscos é fundamental para garantir que a IA seja desenvolvida e utilizada de maneira segura, ética e conforme os valores e direitos fundamentais da União Europeia, proporcionando um equilíbrio entre inovação tecnológica e proteção dos cidadãos.

Quanto a avaliação de riscos associados a utilizações específicas da IA, o regulamento propõe quatro níveis de risco e estabelece regras diferentes nessa conformidade. A Figura 1 apresenta, apresenta um diagrama que demonstra uma pirâmide invertida composta por quatro tipos de sistemas de IA correspondentes a diferentes níveis de risco.

Figura 1. Diagrama exemplificativo dos riscos



Fonte: Conselho da União Europeia

Na base da pirâmide é possível observar a título exemplificativo a utilização de IA com risco mínimo em videogames e filtros de spam, de acordo com o documento, nesse caso a maioria dos sistemas de inteligência artificial (IA) não apresenta riscos significativos e, por isso, poderão continuar a ser utilizados sem restrições ou regulamentações específicas impostas pelo Regulamento de IA.

No segundo nível em *chatbots*, os sistemas de IA se envolveriam no que o regulamento classifica como riscos limitados, e estariam, portanto, sujeitos a obrigações de transparência reduzidas, como a obrigação de informar quando o conteúdo foi gerado por IA permitindo assim que os usuários façam escolhas informadas sobre o uso posterior.

No terceiro piso no uso em transporte, para marcação de exames, recrutamento, concessão de empréstimos, o regulamento pondera que uma ampla gama de sistemas de IA classificados como de risco elevado poderá ser autorizada, desde que cumpram um conjunto rigoroso de requisitos e condições para acessar o mercado da União Europeia.

Em último nível no topo invertido e demonstrando o maior risco, encontra-se a classificação social e reconhecimento facial, classificados como riscos inaceitáveis e, portanto, são proibidos na União Europeia.

São consideradas pelo regulamento inaceitáveis práticas como manipulação cognitiva-comportamental, policiamento preditivo, reconhecimento de emoções em ambientes de trabalho e educativos, e a classificação social. Além disso, o uso de sistemas de identificação biométrica à distância, como reconhecimento facial, será proibido, com algumas exceções específicas.

Em síntese, se de um lado a IA generativa tem uma ampla gama de aplicações, incluindo criação de conteúdo, assistentes virtuais, criação de novos designs de produtos, geração de material didático personalizado para fins de educação acessível, de outra ela apresenta enormes desafios e considerações éticas.

Além do que já foi aduzindo, garantir que o conteúdo gerado seja de alta qualidade, relevante e verdadeiro, preocupações sobre a propagação de preconceitos existentes nos dados de treinamento e o uso ético de conteúdo gerado por IA e questões sobre a autoria e direitos sobre o conteúdo criado por IA são, neste momento, as principais preocupações.

Ao elencar os riscos, o regulamento demonstra importante preocupação com a classificação de indivíduos e grupos por meio da IA.

Afinal, a classificação pode exacerbar desigualdades existentes e introduzir novas formas de discriminação. Sistemas de IA frequentemente dependem de dados históricos para treinamento, o que pode perpetuar preconceitos e estereótipos se esses dados forem enviesados. Assim, o risco de ampliação de desigualdades é um problema central.

Resta saber se essas preocupações também vão tocar a medicina e o sistema de saúde de todo o planeta.

1.4 Riscos, vantagens e desafios para a medicina

A aplicação da inteligência artificial generativa (IA generativa) na medicina e na saúde já oferece uma série de vantagens e riscos. Já é possível a realização de diagnóstico assistido por IA, o que implica maior precisão, celeridade e acessibilidade.

Com efeito, a IA generativa pode analisar grandes volumes de dados médicos, como exames de imagem e históricos de pacientes, ajudando a identificar padrões e anomalias que podem ser difíceis de detectar por médicos humanos. Outrossim, pode acelerar o processo de diagnóstico, permitindo que os médicos tomem decisões mais rápidas e informadas. Mais que tudo isso, em países com dimensões continentais e com enorme desigualdade social, como no Brasil, a possibilidade de ter robôs de IA auxiliando no diagnóstico primário ou acelerando os atendimentos, mesmo à distância, implicará um acesso à saúde de qualidade por pessoas carentes e atualmente desprovidas de atendimento minimamente digno.

Há de se projetar um novo modelo de medicina sob medida: a IA poderá ajudar a criar planos de tratamento personalizados baseados no perfil genético e histórico médico de cada paciente, aumentando a eficácia e diminuindo os custos dos tratamentos, assim como poderá ser usada para identificar novos compostos químicos e prever suas interações, acelerando o desenvolvimento de novos medicamentos (Bensoussan; Gazagne, 2019).

As faculdades de medicina também passarão por uma revolução, permitindo simulações mais realistas para o treinamento de médicos, em ambientes seguros e controlados, além de permitir a criação de conteúdo educacional adaptado às necessidades individuais dos alunos.

A automação de tarefas administrativas aumentará a produtividade, liberando tempo dos profissionais de saúde para se concentrarem em cuidados ao paciente, assim como a gestão de recursos poderá ser otimizada, fomentando a alocação de recursos hospitalares, como leitos e equipamentos médicos.

Na área pública, os mecanismos de transparência, compliance e controle estarão acessíveis a qualquer cidadão.

Como toda revolução e conforme já acentuado no presente artigo, além de bônus, haverá sérios desafios e ônus que precisam ser sopesados.

A falta de precisão e confiabilidade pode permitir que a IA cometa erros de diagnóstico, especialmente se os dados de treinamento não forem representativos ou contiverem vieses.

Eventual dependência excessiva da IA pode ensejar a diminuição das habilidades humanas dos médicos de fazer julgamentos independentes.

Como todo grande banco de dados, há sérios problemas com dados sensíveis, com privacidade e segurança: potencial vazamento ou uso indevido de dados, tanto por criminosos organizados, quanto por seguradoras e empregadores ávidos por aumento nos lucros exigem enorme cautela.

Considerações éticas, como equidade no acesso e dificuldade em atribuir responsabilidade em casos de erro médico envolvendo IA, também implicarão em enormes desafios para a regulação jurídica.

Ademais, não é preciso dizer que haverá enorme impacto nos empregos, seja pela automação de tarefas administrativas e diagnósticas que levará à redução de postos de trabalho em certas áreas da saúde, seja pela evidente necessidade de requalificação, o que pode ser um desafio para alguns profissionais menos adaptados às novas tecnologias.

Em termos práticos, a radiologia já está mudando com algoritmos de IA generativa que podem ajudar a analisar exames de imagem, como raios-X e ressonâncias magnéticas, identificando tumores e outras anomalias com alta precisão (Pereira, 2023, p. 239). A IA pode gerar hipóteses sobre a função de genes e prever mutações que podem causar doenças, ajudando na pesquisa genética e na medicina personalizada, implicando um enorme avanço no biodireito.

De igual modo, a telemedicina, com assistentes virtuais baseados em IA podem fornecer consultas iniciais e triagem de sintomas, direcionando pacientes para os cuidados adequados com base em seus relatos.

Como se vê, a implementação da IA generativa na medicina e na saúde promete transformar muitos aspectos da prática médica e da gestão de saúde, oferecendo grandes benefícios, mas também requerendo cuidadosa consideração dos riscos e desafios envolvidos. Daí a preocupação europeia na regulação que, com o tempo, pode se mostrar açodada e inócua, seja por obstar o avanço tecnológico, seja porque em um mundo

globalizado nem todos têm o compromisso com uma regulação jurídica comum de princípios.

2 Limites da regulação da IA

A regulamentação de sistemas de IA na medicina é um desafio complexo, pois muitas dessas tecnologias lidam com dados sensíveis e estão classificadas como de alto risco. A regulamentação europeia busca equilibrar segurança e inovação, mas regulamentos excessivamente restritivos podem atrasar a introdução de novas tecnologias essenciais para a saúde. Por exemplo, uma IA projetada para a detecção precoce de câncer pode enfrentar demoras significativas em sua aprovação, o que pode resultar em perda de oportunidades de diagnóstico precoce e, conseqüentemente, em vidas que poderiam ter sido salvas.

Por outro lado, a falta de regulamentação pode levar a problemas sérios, como a perpetuação de vieses discriminatórios. Se um sistema de IA é treinado com dados enviesados - por exemplo, imagens de pele clara predominantes - pode falhar em identificar corretamente condições em pessoas com pele mais escura, levando a diagnósticos incorretos e exacerbando desigualdades.

Para evitar esses problemas, é essencial que a regulamentação garanta que os conjuntos de dados utilizados sejam amplamente representativos e diversificados. Incluir informações de diversos grupos demográficos, como diferentes idades, etnias e condições socioeconômicas, ajuda a criar algoritmos mais precisos e equitativos. Além disso, auditorias regulares dos sistemas de IA são fundamentais para identificar e corrigir quaisquer vieses emergentes, garantindo que todos os pacientes recebam tratamento justo.

A transparência também é crucial. Quando os médicos entendem a lógica por trás dos algoritmos, eles podem avaliar e questionar as recomendações da IA garantindo decisões clínicas mais informadas e justas.

Em síntese, enquanto a regulamentação da IA na medicina é necessária para assegurar segurança e equidade, é importante que essa regulamentação seja equilibrada para não limitar a inovação. Aplicar princípios éticos sólidos e garantir a diversidade dos dados, realizar auditorias constantes e promover a transparência são passos fundamentais para desenvolver sistemas de IA que beneficiem todos os pacientes e mantenham o avanço tecnológico essencial para a saúde.

Considerações finais

A proposta de regulamentar a inteligência artificial na União Europeia contribui significativamente para a construção de um modelo de segurança robusto, especialmente ao classificar os sistemas de IA com base em níveis de risco (mínimo, limitado, alto e inaceitável). Este modelo assegura a proteção dos direitos fundamentais, a segurança e a privacidade dos cidadãos, além de promover uma IA confiável e ética. No entanto, como destacado ao longo deste estudo, a efetividade dessa regulação está condicionada à sua adoção em escala global.

A ausência de uma regulamentação harmonizada a nível mundial pode, de fato, resultar em uma competição desleal entre regiões que adotam padrões rigorosos e aquelas que operam sem restrições similares. Isso cria um cenário em que produtos e serviços estrangeiros não regulamentados têm vantagens competitivas, especialmente em áreas sensíveis como a medicina, onde a inovação tecnológica é crucial. O atraso no desenvolvimento científico no Ocidente, causado pela sobrecarga regulatória, é uma preocupação legítima. No entanto, essa preocupação deve ser equilibrada com a necessidade de proteger a saúde pública e os direitos humanos.

Ainda que a regulação possa, a curto prazo, impor desafios à inovação, a criação de uma IA segura e transparente é vital para garantir o avanço científico sustentável. Um modelo de regulação adaptado para mitigar os riscos sem comprometer o progresso científico pode representar uma solução intermediária. A União Europeia, ao liderar com essa iniciativa, estabelece um padrão que outras nações podem eventualmente seguir, criando um caminho para uma regulação mais global e harmonizada, minimizando os impactos negativos no desenvolvimento científico e, ao mesmo tempo, protegendo a sociedade dos riscos inerentes à IA.

Agradecimentos: A autora Andressa Felix Lisboa agradece o apoio concedido pela Coordenação de Aperfeiçoamento de Pessoal de Nível Superior – Brasil (CAPES) – Código de Financiamento 001, durante o desenvolvimento deste estudo.

Referências

BENSOUSSAN, Alain; GAZAGNE, Didier, **Droit des Systèmes Autonomes**, Bruylant, 2019, p. 307.

Comissão Europeia, **Direção-Geral das Redes de Comunicação, Conteúdos e Tecnologias, Orientações éticas para uma IA de confiança, Serviço das Publicações**, 2019.

Comissão Europeia. **Proposta de Regulamento do Parlamento Europeu e do Conselho relativo a uma abordagem harmonizada para a inteligência artificial (Regulamento de IA)**. Eur-Lex, 2021. Disponível em: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:52021PC0206>. Acesso em: 27 jul. 2024.

Conselho da União Europeia. **Acordo sobre o Regulamento de Inteligência Artificial: Conselho e Parlamento alcançam um acordo sobre as primeiras regras mundiais para a IA**. Consilium, 9 dez. 2023. Disponível em: <https://www.consilium.europa.eu/pt/press/press-releases/2023/12/09/artificial-intelligence-act-council-and-parliament-strike-a-deal-on-the-first-worldwide-rules-for-ai/>. Acesso em: 27 jul. 2024.

Conselho da União Europeia. **Proposta de Regulamento do Parlamento Europeu e do Conselho relativo a uma abordagem harmonizada para a inteligência artificial (Regulamento de IA)**. 2022. Disponível em: <https://data.consilium.europa.eu/doc/document/ST-14954-2022-INIT/pt/pdf>. Acesso em: 27 jul. 2024.

Conselho Europeu. **Conclusões do Conselho Europeu de 1 e 2 de outubro de 2020**. Conselho da União Europeia, 2020. Disponível em: <https://www.consilium.europa.eu/media/45928/021020-euco-final-conclusions-pt.pdf>. Acesso em: 27 jul. 2024.

DE MENEZES ACIOLY, Luis Henrique; MENDES, Isabelle Brito Bezerra; NETO, João Araújo Monteiro. **A avaliação de impacto e de resultado regulatório como espectros de política regulatória-sancionatória eficiente em inteligência artificial: reflexões à luz da accountability**. Revista Brasileira de Políticas Públicas, v. 14, n. 1, 2024. Disponível em: <https://www.publicacoesacademicas.uniceub.br/RBPP/article/view/9608>. Acesso em: 28 jul. 2024.

DOI: 10.5281/zenodo.14174188

EUROPA. **Tratado da União Europeia e Tratado sobre o Funcionamento da União Europeia**. Disponível em: <https://eur-lex.europa.eu/legal-content/PT/TXT/HTML/?uri=CELEX:12016ME/TXT#d1e99-13-1>. Acesso em: 27 jul. 2024.

Figura 1: Pirâmide da regulamentação da inteligência artificial. Fonte: Conselho da União Europeia. Disponível em: <https://www.consilium.europa.eu/pt/policies/artificial-intelligence/#0>. Acesso em: 27 jul. 2024.

LAMY, Marcelo. **Metodologia da pesquisa**. Técnicas de investigação, argumentação e redação. 2 ed. revista, atualizada e ampliada. São Paulo: Matrioska, 2020.

MORAES, Alexandre Rocha Almeida de. **Direito Penal do Inimigo: A Terceira Velocidade do direito penal**, Curitiba: Juruá, 2008.

Organização Das Nações Unidas Para A Educação, a ciência e a cultura. **Declaração Universal sobre Bioética e Direitos Humanos**. Paris: UNESCO, 2005. Disponível em: https://unesdoc.unesco.org/ark:/48223/pf0000146180_por. Acesso em: 28 jul. 2024.

ORGANIZAÇÃO PARA A COOPERAÇÃO E DESENVOLVIMENTO ECONÓMICO (OCDE). **Recomendação do Conselho sobre Inteligência Artificial**. Paris: OCDE, 2019. Disponível em: <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>. Acesso em: 28 jul. 2024.

PARLAMENTO EUROPEU. **O que é a inteligência artificial e como funciona**. Parlamento Europeu, 27 ago. 2020. Disponível em: <https://www.europarl.europa.eu/topics/pt/article/20200827STO85804/o-que-e-a-inteligencia-artificial-e-como-funciona>. Acesso em: 27 jul. 2024.

PEREIRA, André Gonçalo Dias. **Inteligência Artificial, Saúde e Direito: considerações jurídicas em torno da medicina de conforto e da medicina transparente**. *Julgar*, (45) (2021), p. 235-262.

POLIDO, F. B. P. **Estado, soberania digital e tecnologias emergentes: interações entre direito internacional, segurança cibernética e inteligência artificial**. *Revista de Ciências do Estado, Belo Horizonte*, v. 9, n. 1, p. 1–30, 2024. DOI: 10.35699/2525-8036.2024.53066. Disponível em: <https://periodicos.ufmg.br/index.php/revce/article/view/e53066>. Acesso em: 28 jul. 2024.

ROSENBLATT, F. (1958). **The Perceptron: A Probabilistic Model for Information Storage and Organization in The Brain**. *Psychological Review*, 65(6), p. 386-408.

RUMELHART, D. E., HINTON, G. E., & WILLIAMS, R. J. (1986). **Learning representations by back-propagating errors**. *Nature*, 323(6088), 533-536

UNIÃO EUROPEIA. **Regulamento do Parlamento Europeu e do Conselho relativo à abordagem europeia à inteligência artificial. PE-24-2024-INIT**. Disponível em: <https://data.consilium.europa.eu/doc/document/PE-24-2024-INIT/pt/pdf>. Acesso em: 28 jul. 2024.

Informação bibliográfica deste texto, conforme a NBR 6023:2018 da Associação Brasileira de Normas Técnicas (ABNT):

MORAES, Alexandre Rocha Almeida de; LISBOA, Andressa Felix. Regulamentação da Inteligência Artificial na União Europeia: estrutura ética, classificação de riscos e possíveis reflexos na medicina. **UNISANTA Law and Social Science**, Vol. 13, N. 2 (jul./dez. 2024), pp. 16-29. ISSN: 2317-1308.

Recebido em 10/08/2024
Aprovado em 20/10/2024



<https://creativecommons.org/licenses/by/4.0/deed.pt-br>

DOI: 10.5281/zenodo.14174188