

UNIVERSIDADE FEDERAL DE CIÊNCIAS DA SAÚDE DE PORTO ALEGRE
CURSO DE GESTÃO EM SAÚDE

Matheus de Matos Medina

FRAMEWORKS REGULATÓRIOS DE INTELIGÊNCIA ARTIFICIAL NA SAÚDE:
UMA ANÁLISE DOCUMENTAL COMPARATIVA

Porto Alegre
2025

Matheus de Matos Medina

**FRAMEWORKS REGULATÓRIOS DE INTELIGÊNCIA ARTIFICIAL NA SAÚDE:
UMA ANÁLISE DOCUMENTAL COMPARATIVA**

Trabalho de Conclusão de Curso a ser apresentado
no Curso de Graduação em Gestão em Saúde,
como requisito obrigatório para obtenção Grau de
Bacharel em Gestão em Saúde pela Universidade
Federal de Ciências da Saúde de Porto Alegre.

Orientadora: Profa. Dra. Melissa Medeiros
Markoski

Porto Alegre
2025

Catálogo na Publicação

Medina, Matheus de Matos
FRAMEWORKS REGULATÓRIOS DE INTELIGÊNCIA ARTIFICIAL NA
SAÚDE: UMA ANÁLISE DOCUMENTAL COMPARATIVA / Matheus de
Matos Medina. -- 2025.
66 p. : 30 cm.

Monografia (trabalho de conclusão de curso) --
Universidade Federal de Ciências da Saúde de Porto
Alegre, Curso de Gestão em Saúde, 2025.

Orientador(a): Profa. Dra. Melissa Medeiros Markoski.

1. inteligência artificial. 2. regulação em saúde. 3.
governança de dados. 4. análise comparativa. I. Título.

Sistema de Geração de Ficha Catalográfica da UFCSPA com os dados
fornecidos pelo(a) autor(a).

“Nenhum esforço faz sentido, se você não acredita em si mesmo”.

Masashi Kishimoto (Maito Gai)

AGRADECIMENTOS

Agradeço, primeiramente, à minha orientadora, Prof.^a Mel, pela parceria fundamental durante os últimos dois anos. Sou imensamente grato(a) por todo o suporte oferecido durante nosso projeto de iniciação à docência e na condução deste Trabalho de Conclusão de Curso.

À equipe de qualificação e educação continuada do Hospital de Clínicas, que me oportunizou um imenso crescimento profissional e proporcionou dias mais tranquilos e enriquecedores durante a graduação.

Ao projeto de extensão "Cuidando da Farmácia Caseira", que se tornou um refúgio e trouxe momentos de paz em toda a minha trajetória, além de me permitir levar o conhecimento para além dos muros da universidade.

A todos do laboratório LAIT, mas especialmente às Dras. Claudinha Paiva e Cristina Bonorino, por sempre acreditarem em meu potencial e por todo o suporte na reta final.

À querida professora Mellina, por ter me concedido a valiosa oportunidade de ingressar na iniciação científica, um passo decisivo em minha formação.

À minha família, em especial à minha querida mãe e minha tia, pelo incentivo incondicional desde o início e por serem a base que me permitiu concluir esta graduação.

Este trabalho de pesquisa só foi possível através do apoio e suporte do meu amor, Jordana K. Por todo o amor, paciência e fortalecimento nos momentos adversos, dedico esta pesquisa a ela.

RESUMO

A crescente integração da inteligência artificial no setor da saúde impõe novos desafios regulatórios. No Brasil, este cenário é caracterizado por uma soberania dupla, com a sobreposição de normas da ANVISA e da LGPD, e pela tramitação do PL nº 2.338/2023. Este trabalho analisou comparativamente os frameworks regulatórios de IA na saúde dos Estados Unidos, União Europeia e Singapura, identificando tendências e lições estratégicas aplicáveis ao Brasil. Foi conduzida pesquisa documental comparativa de 18 documentos oficiais baseada em cinco categorias analíticas derivadas dos princípios de IA confiável da OECD: transparência e explicabilidade, robustez e segurança, responsabilidade, justiça e privacidade, e governança do ciclo de vida do algoritmo. Os resultados revelaram dois arquétipos regulatórios: o garantista-prescritivo, caracterizado por hard law detalhado e supervisão estatal direta, liderado pela União Europeia e seguido pelo Brasil; e o pragmático-processual, caracterizado por soft law orientador e governança organizacional, liderado pelos Estados Unidos e Singapura. A análise posicionou o Brasil em trajetória de convergência com o modelo europeu garantista-prescritivo, refletindo valores constitucionais de proteção de direitos fundamentais. O trabalho contribui para os Objetivos de Desenvolvimento Sustentável 3, 10 e 16, oferecendo referencial para gestores, formuladores de políticas e reguladores no momento crítico de definição do marco regulatório de IA no Brasil.

Palavras-chave: inteligência artificial; regulação em saúde; governança de dados; análise comparativa.

ABSTRACT

The growing integration of artificial intelligence in healthcare poses new regulatory challenges. In Brazil, this scenario is characterized by overlapping regulatory authorities, with simultaneous requirements from ANVISA and LGPD, and the pending bill 2338/2023. This study comparatively analyzed AI regulatory frameworks in healthcare from the United States, European Union, and Singapore, identifying trends and strategic lessons applicable to Brazil. Comparative documentary research was conducted analyzing 18 official documents based on five analytical categories derived from OECD's trustworthy AI principles: transparency and explainability, robustness and security, accountability, fairness and privacy, and algorithmic life cycle governance. Results revealed two regulatory archetypes: the rights-based prescriptive model, characterized by detailed hard law and direct state supervision, led by the European Union and followed by Brazil; and the pragmatic-procedural model, characterized by guiding soft law and organizational governance, led by the United States and Singapore. The analysis positioned Brazil on a convergence trajectory with the European rights-based prescriptive model, reflecting constitutional values of fundamental rights protection. This work contributes to UN Sustainable Development Goals 3, 10, and 16, offering a strategic framework for health managers, policymakers, and regulators at the critical moment of defining Brazil's AI regulatory framework.

Keywords: artificial intelligence; health regulation; data governance; comparative analysis.

LISTA DE ABREVIATURAS

AI Act - *Artificial Intelligence Act* (Regulamento de Inteligência Artificial da União Europeia)

AIHGle - *Artificial Intelligence in Healthcare Guidelines* (Diretrizes de IA em Saúde de Singapura)

ANPD - Agência Nacional de Proteção de Dados

ANVISA - Agência Nacional de Vigilância Sanitária

CCPA - *California Consumer Privacy Act* (Lei de Privacidade do Consumidor da Califórnia)

CONEP - Comissão Nacional de Ética em Pesquisa

DPIA - Avaliações de Impacto sobre a Proteção de Dados

EUA - Estados Unidos da América

FDA - *Food and Drug Administration* (Agência Federal de Alimentos e Medicamentos dos EUA)

GDPR - *General Data Protection Regulation* (Regulamento Geral sobre a Proteção de Dados)

HHS - *Department of Health and Human Services* (Departamento de Saúde e Serviços Humanos dos EUA)

HIPAA - *Health Insurance Portability and Accountability Act* (Lei de Portabilidade e Responsabilidade de Seguros de Saúde dos EUA)

IA - Inteligência Artificial

IEC - *International Electrotechnical Commission* (Comissão Eletrotécnica Internacional)

IMDA - *Infocomm Media Development Authority* (Autoridade de Desenvolvimento de Mídia e Infocomm de Singapura)

ISO - *International Organization for Standardization* (Organização Internacional de Normalização)

LGPD - Lei Geral de Proteção de Dados Pessoais (Lei nº 13.709/2018)

MDR - *Medical Device Regulation*

MHRA - *Medicines and Healthcare products Regulatory Agency* (Agência Reguladora de Medicamentos e Produtos de Saúde do Reino Unido)

NIST - *National Institute of Standards and Technology* (Instituto Nacional de Padrões e Tecnologia dos EUA)

OECD - *Organisation for Economic Co-operation and Development* (Organização para a Cooperação e Desenvolvimento Econômico)

PCCP - *Predetermined Change Control Plan* (Plano de Controle de Mudanças Pré-determinado)

PDPA - *Personal Data Protection Act* (Lei de Proteção de Dados Pessoais de Singapura)

PDPC - *Personal Data Protection Commission* (Comissão de Proteção de Dados Pessoais de Singapura)

PETs - *Privacy-Enhancing Technologies* (Tecnologias de Melhoria da Privacidade)

PL - Projeto de Lei

RDC - Resolução da Diretoria Colegiada

SaMD - *Software as a Medical Device* (Software como Dispositivo Médico)

SIA - Sistema Nacional de Regulação e Governança de Inteligência Artificial

SLA - *Service Level Agreement* (Acordo de Nível de Serviço)

SUS - Sistema Único de Saúde

UE - União Europeia

SUMÁRIO

RESUMO	6
ABSTRACT	7
LISTA DE ABREVIATURAS	8
1. INTRODUÇÃO	12
2. OBJETIVOS	13
2.1. OBJETIVO GERAL	13
2.2. OBJETIVOS ESPECÍFICOS	13
3. JUSTIFICATIVA	14
3.1 PROBLEMA DE PESQUISA	15
4. REFERENCIAL TEÓRICO	16
4.1 O CONCEITO DE FRAMEWORK E SUA APLICAÇÃO NA GOVERNANÇA DE IA EM SAÚDE	16
4.2 A IA E A GESTÃO NOS SERVIÇOS DE SAÚDE	17
4.3 O PARADOXO DA IA EM SAÚDE	18
4.4 DESAFIOS CENTRAIS PARA A REGULAÇÃO DE IA EM SAÚDE	19
4.5 CICLO DE VIDA DA IA	21
4.6 OS PRINCÍPIOS DE IA CONFIÁVEL DA OECD	21
4.7 HARD LAW E SOFT LAW NA GOVERNANÇA DE IA	23
5. METODOLOGIA	24
5.1 DELINEAMENTO DA PESQUISA	24
5.2 SELEÇÃO DOS CASOS PARA ANÁLISE COMPARATIVA	25
5.3 FONTES DE DADOS E DELIMITAÇÃO DO CORPUS	25
5.4 PROCEDIMENTOS DE ANÁLISE	27
6. RESULTADOS	28
6.1 ANÁLISE DESCRITIVA DOS FRAMEWORKS REGULATÓRIOS	29
6.1.1 O FRAMEWORK REGULATÓRIO DO BRASIL	29
6.1.1.1 ARQUITETURA DO FRAMEWORK E NATUREZA DOS INSTRUMENTOS	29
6.1.1.2 ANÁLISE CATEGORIAL	30
6.1.2 O FRAMEWORK REGULATÓRIO DA UE	31
6.1.2.1 ARQUITETURA DO FRAMEWORK E NATUREZA DOS INSTRUMENTOS	31
6.1.2.2 ANÁLISE CATEGORIAL	32

6.1.3 O FRAMEWORK REGULATÓRIO DO EUA.....	34
6.1.3.1 ARQUITETURA DO FRAMEWORK E NATUREZA DOS INSTRUMENTOS	34
6.1.3.2 ANÁLISE CATEGORIAL	34
6.1.4 O FRAMEWORK REGULATÓRIO DE SINGAPURA	36
6.1.4.1 ARQUITETURA DO FRAMEWORK E NATUREZA DOS INSTRUMENTOS	36
6.1.4.2 ANÁLISE CATEGORIAL	37
7. ANÁLISE COMPARATIVA	38
7.1 MATRIZ DE ANÁLISE COMPARATIVA	39
7.2 DISCUSSÃO DAS CONVERGÊNCIAS E DIVERGÊNCIAS ESTRATÉGICAS.....	41
7.2.1 TRANSPARÊNCIA E EXPLICABILIDADE.....	41
7.2.2 ROBUSTEZ, SEGURANÇA E PROTEÇÃO	43
7.2.3 RESPONSABILIDADE (ACCOUNTABILITY).....	44
7.2.4 JUSTIÇA E PRIVACIDADE	46
7.2.5 GOVERNANÇA DO CICLO DE VIDA DO ALGORITMO	48
7.3 SÍNTESE DA ANÁLISE CATEGORIAL	50
8. DISCUSSÃO	53
8.1 SÍNTESE DOS ACHADOS DA PESQUISA.....	53
8.2 O POSICIONAMENTO REGULATÓRIO DO BRASIL	55
8.3 IMPLICAÇÕES PARA A GESTÃO EM SAÚDE.....	56
8.4 RECOMENDAÇÕES	57
8.5 LIMITAÇÕES DO ESTUDO	58
8.6 PESQUISAS FUTURAS	59
9. CONCLUSÃO	60
10. REFERÊNCIAS	62

1. INTRODUÇÃO

A crescente integração da IA no setor da saúde representa uma das transformações mais disruptivas e promissoras da medicina moderna. Com o potencial de aprimorar a acurácia diagnóstica, otimizar processos de gestão e personalizar tratamentos, as ferramentas de IA estão deixando de ser uma promessa futurista para se tornarem uma realidade operacional em instituições de saúde ao redor do mundo. Para o gestor em saúde, compreender e adotar essas tecnologias tornou-se um imperativo estratégico, essencial para a melhoria da qualidade assistencial e para a sustentabilidade das organizações frente a um cenário de crescente complexidade e demanda por eficiência.

Contudo, este avanço tecnológico avança em um ritmo muito superior ao do desenvolvimento de seus marcos legais. A implementação de sistemas de IA, especialmente aqueles que apoiam decisões clínicas de alto impacto e processam dados sensíveis de pacientes, introduz um panorama de riscos jurídicos, éticos e operacionais sem precedentes. No Brasil, as instituições de saúde se encontram em um complexo regulatório, submetidas simultaneamente às normas de dispositivos médicos da ANVISA e aos mandatos de proteção de dados da LGPD.

Este cenário de soberania dupla é agravado pela iminente aprovação de uma legislação específica para IA, o PL nº 2.338/2023, que promete redefinir as regras de governança e responsabilidade. Essa conjuntura cria incertezas tanto para formuladores de políticas públicas quanto para gestores de instituições de saúde, uma vez que decisões sobre regulação e implementação de tecnologias de IA são tomadas sem referências claras sobre as melhores práticas internacionais, aumentando os riscos jurídicos, éticos e operacionais para todos os atores envolvidos no ecossistema de saúde.

A complexidade regulatória é amplificada pela coexistência de diferentes instrumentos normativos que variam em sua força vinculante. Alguns países adotam abordagens de *hard law*, isto é, regulações juridicamente vinculantes com mecanismos de fiscalização e aplicação coercitiva, enquanto outros privilegiam instrumentos de *soft law*, como diretrizes, princípios e padrões de caráter voluntário ou orientativo. Essa diversidade de abordagens normativas reflete diferentes estratégias para equilibrar segurança, inovação e proteção de direitos no contexto da IA em saúde, aspectos que serão explorados na análise comparativa dos *frameworks* regulatórios apresentados neste trabalho.

Apesar do intenso debate global, observa-se uma lacuna na produção acadêmica nacional: a ausência de uma análise comparativa sistemática que coloque o emergente modelo brasileiro em diálogo com os principais paradigmas regulatórios internacionais. Modelos como

o da FDA nos EUA, focado na inovação adaptativa, e o do AI Act na UE, centrado nos direitos fundamentais, oferecem lições valiosas que ainda não foram devidamente traduzidas para a realidade do cenário regulatório e da governança de IA no Brasil.

Diante desta lacuna, o presente trabalho se propõe a realizar uma análise documental comparativa dos frameworks regulatórios de IA na saúde dos EUA, UE e Singapura, contrastando-os com o cenário brasileiro. Essas três jurisdições foram selecionadas por representarem arquétipos regulatórios distintos como regulação baseada em direitos (UE), inovação adaptativa (EUA) e governança colaborativa (Singapura), oferecendo um espectro amplo de abordagens aplicáveis ao contexto brasileiro. O objetivo é extrair lições estratégicas e identificar diretrizes aplicáveis ao aprimoramento do cenário regulatório brasileiro e à governança de IA nas instituições de saúde.

Para tanto, este trabalho de conclusão de curso está estruturado da seguinte forma. Após esta introdução, o Capítulo 2 apresenta os objetivos da pesquisa, e o Capítulo 3, sua justificativa e problema de pesquisa. O Capítulo 4 estabelece o referencial teórico, abordando os conceitos fundamentais de IA em saúde e os desafios centrais para sua regulação. O Capítulo 5 descreve a metodologia de pesquisa documental comparativa adotada. Os Capítulos 6 e 7 apresentam os resultados da análise: o primeiro caracteriza os frameworks regulatórios de cada jurisdição estudada; o segundo realiza a análise comparativa, identificando convergências, divergências e lições estratégicas. Por fim, o Capítulo 8 sintetiza as conclusões, discutindo o posicionamento regulatório do Brasil, as implicações para a governança de IA nas instituições de saúde e apresentando recomendações para diferentes atores do ecossistema.

2. OBJETIVOS

2.1. OBJETIVO GERAL

Analisar comparativamente os frameworks regulatórios de IA na saúde dos EUA, UE e Singapura, a fim de identificar as principais tendências, desafios e lições estratégicas aplicáveis ao cenário regulatório brasileiro e à governança de IA nas instituições de saúde.

2.2. OBJETIVOS ESPECÍFICOS

- Caracterizar os modelos regulatórios de referência para IA em saúde dos EUA, UE e Singapura, identificando seus arquétipos e mecanismos centrais.

- Caracterizar o arcabouço legal brasileiro aplicável à IA na saúde, abrangendo as normativas da ANVISA (RDC 751/2022), a LGPD e o PL nº 2.338/2023.
- Comparar os diferentes frameworks regulatórios com base em categorias analíticas pré-definidas, identificando convergências, divergências e lacunas entre os modelos internacionais e o brasileiro.
- Discutir as implicações estratégicas da análise comparativa para o contexto brasileiro, destacando os principais desafios para a governança de IA nas instituições de saúde.

3. JUSTIFICATIVA

A crescente integração da IA no ecossistema de saúde representa uma das mais significativas transformações do setor. Conforme destacam Palaniappan, Lin e Vogel (2024), as aplicações de IA são vastas e esperadas para auxiliar, automatizar e aumentar diversos serviços de saúde, com o objetivo final de melhorar os resultados e a satisfação do paciente, a saúde da população e o bem-estar dos profissionais. Para o gestor em saúde, a adoção dessas tecnologias não é mais uma questão de "se", mas de "como" e "quando", tornando-se um imperativo para a competitividade e a melhoria da qualidade assistencial.

Contudo, este avanço tecnológico traz consigo uma complexidade regulatória sem precedentes. A implementação de sistemas de IA, especialmente aqueles que lidam com dados sensíveis de pacientes e apoiam decisões de vida ou morte, expõe as instituições de saúde a riscos significativos de ordem legal, ética e reputacional. O desafio é amplificado no Brasil por um cenário de "soberania dupla", onde um mesmo software pode estar sujeito simultaneamente às rigorosas normas de dispositivos médicos da ANVISA e aos mandatos de proteção de dados e direito à explicação da LGPD.

Esta dualidade cria uma zona de incerteza para o gestor, que precisa garantir a conformidade em múltiplas frentes, muitas vezes com requisitos distintos. A situação se torna ainda mais complexa com a tramitação do PL nº 2.338/2023, que propõe um novo marco legal para a IA no Brasil. Este futuro cenário exigirá das instituições de saúde não apenas conhecimento técnico, mas uma profunda literacia regulatória para tomar decisões estratégicas sobre aquisição, implementação e governança dessas tecnologias.

É neste contexto que a presente pesquisa se justifica. Apesar do crescente debate internacional sobre a regulação da IA, observa-se uma lacuna na produção acadêmica nacional: a ausência de uma análise comparativa sistemática que coloque o modelo brasileiro em diálogo

com os principais frameworks globais. Sem essa análise, o gestor de saúde brasileiro fica sem referências claras para antecipar tendências e adotar as melhores práticas internacionais.

Portanto, a contribuição deste trabalho é preencher essa lacuna ao realizar uma análise documental comparativa detalhada. Ao mapear as convergências, divergências e inovações dos principais modelos regulatórios, esta pesquisa visa fornecer um quadro analítico que sirva como um guia estratégico. O estudo não se limita a descrever as leis ou atos regulatórios, mas busca extrair lições práticas e discutir suas implicações para a governança de IA nas instituições de saúde brasileiras.

A seleção dos EUA, UE e Singapura como modelos de referência para esta análise comparativa fundamenta-se em critérios que consideram diferentes abordagens regulatórias existentes no cenário internacional. Estes três casos foram escolhidos por representarem arquétipos distintos e filosofias de governança diferentes para IA em saúde. A UE destaca-se por sua regulação baseada em direitos fundamentais, tendo desenvolvido o primeiro marco legal horizontal específico para IA (*AI Act*) com obrigações vinculantes detalhadas, oferecendo exemplos de como converter princípios éticos em requisitos técnicos verificáveis. Os EUA adotam abordagem voltada para inovação, utilizando principalmente diretrizes voluntárias através de documentos técnicos como o NIST AI RMF e mecanismos setoriais específicos como o *Predetermined Change Control Plan* da FDA, demonstrando formas de combinar flexibilidade regulatória com supervisão de segurança. Singapura representa um modelo de governança colaborativa, onde o Estado atua como facilitador da autorregulação através de guias técnicos e parcerias público-privadas, estabelecendo-se como centro de desenvolvimento de IA na Ásia.

A análise desses três modelos permite ao Brasil, que está em processo de definição de seu marco regulatório, compreender as diferentes escolhas regulatórias e suas implicações para proteção de direitos, estímulo à inovação e organização institucional. O estudo comparativo busca identificar elementos aplicáveis ao contexto brasileiro, considerando características como o SUS e a necessidade de reduzir desigualdades no acesso à tecnologia em saúde. É nesse contexto que se formula o problema central desta investigação.

3.1 PROBLEMA DE PESQUISA

Diante do dinâmico cenário regulatório apresentado, e considerando uma possível lacuna na literatura nacional que ofereça um guia comparativo para a tomada de decisão,

emerge a necessidade de uma investigação que traduza as experiências internacionais em conhecimento estratégico para a realidade brasileira. A simples adoção de tecnologias de IA, sem uma compreensão profunda de suas implicações legais e de governança, representa um risco significativo para as instituições de saúde e seus gestores. Nesse sentido, a presente pesquisa busca responder à seguinte questão central:

A partir da análise comparativa dos frameworks regulatórios de IA em saúde dos EUA, UE e Singapura, quais são as principais diretrizes e desafios para o aprimoramento do cenário regulatório brasileiro e para a governança eficaz de IA nas instituições de saúde?

4. REFERENCIAL TEÓRICO

Este capítulo estabelece a fundamentação conceitual para a análise dos frameworks regulatórios de IA na saúde. Serão abordados os principais usos da IA no setor, os desafios éticos e operacionais que sua implementação acarreta, e o conceito de governança de IA como uma resposta estratégica.

4.1 O CONCEITO DE FRAMEWORK E SUA APLICAÇÃO NA GOVERNANÇA DE IA EM SAÚDE

O termo *framework*, amplamente utilizado em contextos de governança e regulação, refere-se a uma estrutura sistemática composta por princípios, diretrizes, processos e ferramentas que orientam a abordagem de questões complexas (OECD, 2019). No contexto da governança de inteligência artificial em saúde, a OMS estabelece que um framework regulatório compreende um conjunto integrado de princípios éticos, leis, padrões e diretrizes que orientam o desenvolvimento e implementação responsável da tecnologia (WHO, 2021).

Diferentemente de regulações prescritivas isoladas, um framework integra múltiplos componentes regulatórios em uma estrutura coerente, característica que se mostra particularmente adequada para contextos de rápida evolução tecnológica.

A importância dos *frameworks* regulatórios reside em sua capacidade de operacionalizar princípios abstratos, traduzindo valores éticos e requisitos legais em orientações práticas. Na governança de IA, diferentes tipos de *frameworks* cumprem funções complementares: estruturas legais estabelecem proteções para privacidade e consentimento; marcos de direitos humanos fornecem bases para responsabilização; e arcabouços de governança definem regras de engajamento entre setores público e privado. Contudo, análises recentes indicam que as

estruturas regulatórias existentes permanecem fragmentadas e muitos países ainda carecem de mecanismos adequados para governar IA em saúde (OMS, 2021). Assim, o desenvolvimento de frameworks robustos emerge como ferramenta fundamental para assegurar que a inteligência artificial seja implementada de forma responsável e benéfica no setor saúde.

4.2 A IA E A GESTÃO NOS SERVIÇOS DE SAÚDE

A inteligência artificial pode ser compreendida como a capacidade de sistemas computacionais de executar tarefas que normalmente exigiria inteligência humana, tais como reconhecimento de padrões, tomada de decisões e aprendizado a partir de experiências (JIANG et al., 2017). Embora o termo abranja diferentes abordagens técnicas, a IA contemporânea em saúde é predominantemente baseada em aprendizado de máquina, técnica que permite aos algoritmos aprenderem padrões a partir de grandes volumes de dados sem serem explicitamente programados para cada situação específica. As redes neurais profundas, inspiradas no funcionamento do cérebro humano, representam a forma mais avançada desse aprendizado, sendo capazes de identificar relações complexas em dados clínicos, imagens médicas e registros de pacientes. É essa capacidade de aprender e se adaptar a partir de dados que distingue a IA moderna de sistemas computacionais tradicionais e que fundamenta tanto suas promessas quanto seus riscos regulatórios.

Na prática clínica, sistemas de IA baseados em aprendizado profundo têm demonstrado capacidade de auxiliar profissionais de saúde em tarefas diagnósticas complexas. O estudo de Esteva et al. (2017) exemplifica essa capacidade ao desenvolver uma rede neural convolucional para classificação de câncer de pele treinada com mais de 129.000 imagens clínicas. Quando comparado ao desempenho de 21 dermatologistas certificados, o sistema alcançou nível de competência equivalente na identificação de carcinomas de células basais e melanomas malignos, evidenciando o potencial dessas tecnologias para ampliar o acesso a diagnósticos especializados, especialmente em regiões com escassez de especialistas. Além de diagnósticos por imagem, a IA é aplicada na interpretação de exames laboratoriais, predição de riscos clínicos, recomendação de tratamentos e triagem de pacientes, transformando a prática médica em diversas especialidades.

Para além das aplicações clínicas, a IA está remodelando a gestão e a administração dos serviços de saúde. Zhang (2024) argumentou que sistemas de IA possuem potencial significativo para melhorar a eficiência de gestão hospitalar, otimizando a alocação de recursos, prevendo demanda por serviços, automatizando processos administrativos e apoiando decisões

gerenciais baseadas em dados. Algoritmos preditivos são utilizados para antecipar picos de demanda em emergências, otimizar escalas de profissionais e reduzir desperdícios operacionais. Essa expansão acelerada da IA tanto em contextos clínicos quanto administrativos, entretanto, não ocorre sem desafios substanciais que exigem atenção regulatória cuidadosa, conforme será explorado na próxima seção.

4.3 O PARADOXO DA IA EM SAÚDE

A implementação da IA na saúde revela um paradoxo: quanto maior o potencial transformador da tecnologia, mais complexos tornam-se os desafios éticos e regulatórios que condicionam sua realização segura. Gerke, Minssen e Cohen (2020) destacaram que embora a IA ofereça um potencial de melhorias para o sistema de saúde, a plena concretização desses benefícios está intrinsecamente condicionada à superação de riscos substanciais inerentes à sua aplicação. Essa tensão entre promessa e perigo se manifesta de no setor da saúde, onde decisões algorítmicas podem impactar diretamente a vida e a saúde dos pacientes, tornando a urgência por marcos regulatórios robustos ainda mais evidente.

Vollmer et al. (2020) alertam que a aplicação da IA é frequentemente limitada por carência de transparência, replicabilidade e demonstrações claras de eficácia clínica no mundo real. Essa lacuna entre a promessa tecnológica e a evidência robusta gera o risco de permitir a tradução de algoritmos ineficazes ou mesmo prejudiciais da bancada do computador para o leito do paciente. O viés algorítmico exemplifica de forma contundente esse risco: Obermeyer et al. (2019) documentaram como um algoritmo amplamente utilizado subestimava sistematicamente a gravidade das doenças em pacientes negros ao empregar custos de saúde como indicador indireto para necessidades de saúde, perpetuando e amplificando iniquidades estruturais ao reduzir o encaminhamento de populações vulneráveis para programas de cuidado especializado.

A dependência de grandes volumes de dados para treinamento intensifica adicionalmente os riscos à privacidade dos pacientes. Price II e Cohen (2019) argumentaram que dados de saúde são particularmente sensíveis, e seu acúmulo em larga escala permite o conhecimento de informações íntimas por terceiros, seja através de divulgação inadvertida ou ataques cibernéticos, levantando preocupações sobre segurança e uso apropriado das informações. Esses riscos de opacidade algorítmica, viés discriminatório, responsabilidade difusa e vulnerabilidade de dados exigem respostas regulatórias específicas e robustas, cujos contornos conceituais serão explorados nas seções subsequentes.

4.4 DESAFIOS CENTRAIS PARA A REGULAÇÃO DE IA EM SAÚDE

O avanço das redes neurais profundas na ciência introduz o problema da caixa-preta, onde modelos de inteligência artificial funcionam de maneira tão opaca quanto o cérebro humano, difundindo o conhecimento adquirido de uma forma extremamente difícil de decifrar (CASTELVECCHI, 2016). Embora a máquina possa oferecer respostas precisas, o conhecimento fica incorporado na rede, não permitindo que humanos compreendam o processo decisório. Um cientista não se satisfaz em apenas distinguir entre gatos e cães; ele quer ser capaz de explicar as características que diferenciam essas categorias, algo que a natureza não interpretável dessas caixas-pretas impede. No contexto clínico, essa opacidade torna-se ainda mais problemática: quando um sistema de IA recomenda um tratamento específico ou identifica uma lesão como maligna, médicos e pacientes precisam compreender as razões por trás dessa conclusão para validar a recomendação e tomar decisões informadas.

A opacidade algorítmica transcende o campo técnico para se tornar um significativo risco jurídico e operacional. Este risco é materializado por mandatos legais como o Artigo 20 da LGPD, que assegura ao titular o direito à revisão de decisões tomadas unicamente com base em tratamento automatizado (BRASIL, 2018). Conforme analisou Bioni (2019), a dificuldade em prover informações claras e adequadas a respeito dos critérios e dos procedimentos utilizados para a decisão automatizada representa um dos maiores desafios para a conformidade regulatória, transformando a complexidade algorítmica em um obstáculo direto ao exercício de um direito pelo cidadão. Para gestores de saúde, esse desafio se materializa no momento da aquisição de tecnologia: é necessário avaliar não apenas a acurácia do sistema, mas sua capacidade de fornecer explicações compreensíveis sobre suas decisões, equilibrando desempenho técnico com transparência operacional e conformidade legal.

A pesquisa em Inteligência Artificial Explicável, conhecida pela sigla xAI, dedica-se a superar o problema da caixa-preta através do desenvolvimento de técnicas que tornam os sistemas algorítmicos mais compreensíveis (ARRIETA et al., 2020). Na prática clínica, isso pode significar, por exemplo, que um sistema de diagnóstico de câncer de pele não apenas classifique uma lesão como maligna, mas também indique quais características visuais específicas (assimetria, irregularidade de bordas, variação de cor) influenciaram essa classificação, permitindo que o dermatologista valide ou questione a recomendação com base em seu conhecimento clínico.

O viés algorítmico emerge como o segundo desafio fundamental para a governança de IA, pois os modelos de aprendizado de máquina aprendem a partir de dados que refletem vieses

humanos existentes ou iniquidades em nível de sistema. Consequentemente, ao serem treinados com um reflexo do mundo como ele é, esses sistemas podem não apenas perpetuar, mas até mesmo amplificar as desigualdades já presentes no cuidado em saúde, resultando em disparidades que podem levar a danos aos pacientes (WIENS et al., 2019). O caso documentado por Obermeyer et al. (2019), discutido anteriormente na seção 4.2, exemplifica como um algoritmo amplamente utilizado subestimou a gravidade das doenças em pacientes negros ao empregar custos de saúde como proxy para necessidades de saúde, perpetuando desigualdades estruturais no acesso ao cuidado. Esse exemplo demonstra que o viés não decorre necessariamente de intenção discriminatória dos desenvolvedores, mas da escolha inadequada de variáveis que refletem e amplificam iniquidades sociais preexistentes.

A sub-representação de dados demográficos constitui uma fonte crítica adicional de viés algorítmico. Conforme demonstraram Cirillo et al. (2020), os dados usados para treinar sistemas de IA frequentemente sub-representam as mulheres, o que pode fazer com que os sistemas de IA tenham um desempenho inferior para mulheres em comparação com homens, levando a um aumento das desigualdades de gênero na saúde ao ser aplicada na prática clínica. Para gestores de saúde, isso impõe a necessidade de avaliar criticamente os conjuntos de dados utilizados no treinamento dos sistemas de IA que pretendem adquirir, questionando se esses dados são suficientemente representativos das populações que serão atendidas pela instituição. Um algoritmo treinado predominantemente com imagens de pele clara, por exemplo, pode apresentar desempenho inadequado na detecção de melanoma em pacientes de pele escura, criando disparidades no acesso a diagnósticos precisos.

O terceiro desafio conceitual reside na definição de responsabilidade por danos decorrentes de erros de IA na saúde. Price II, Gerke e Cohen (2019) abordaram essa questão através do conceito de um ecossistema maior de responsabilidade, onde a falha não reside em um único agente. Nesse modelo, a responsabilidade do médico é vista como apenas uma peça de uma cadeia que também inclui os fabricantes da IA médica e os sistemas hospitalares que compram e implementam a tecnologia. Quando um sistema de IA recomenda um tratamento inadequado que resulta em dano ao paciente, estabelecer a responsabilidade exige determinar se o erro se originou de dados de treinamento inadequados (responsabilidade do desenvolvedor), de implementação inadequada ou falta de supervisão humana (responsabilidade da instituição), ou de confiança excessiva na recomendação algorítmica sem validação clínica apropriada (responsabilidade do profissional). Essa distribuição difusa de responsabilidade cria incerteza jurídica e operacional, exigindo que gestores de saúde

estabeleçam claramente as responsabilidades contratuais com fornecedores e definam protocolos internos de supervisão humana sobre decisões algorítmicas.

4.5 CICLO DE VIDA DA IA

Os desafios de transparência, viés e responsabilidade, detalhados anteriormente, não são eventos isolados, mas se manifestam ao longo de todo o ciclo de vida do sistema de IA. Este conceito, central para as regulações modernas como a da OECD (2019) e o PL 2.338/2023, compreende todas as fases do sistema, desde o planejamento até sua desativação. A governança eficaz exige uma mudança de mentalidade: da avaliação de um produto estático para o gerenciamento de um processo dinâmico e contínuo.

A aplicação prática dessa governança pode ser visualizada nas etapas-chave do ciclo de vida, onde os riscos conceituais se materializam. Na fase de identificação do problema e coleta de dados, os riscos de viés e privacidade são mais latentes, exigindo conformidade com a LGPD e uma curadoria cuidadosa dos dados de treinamento para garantir a justiça algorítmica. Por fim, após a implementação e durante o monitoramento e manutenção, a governança foca na supervisão humana e na gestão da evolução do sistema, onde o desafio do modelo "estático vs. dinâmico" se materializa e a obrigação de manter a responsabilidade (*accountability*) exige trilhas de auditoria robustas.

Um exemplo prático dessa abordagem de monitoramento contínuo é o modelo PCCP implementado pela FDA (2023). Esse *framework* reconhece que algoritmos de aprendizado contínuo evoluem após sua aprovação inicial, estabelecendo um protocolo onde fabricantes devem especificar antecipadamente as modificações planejadas, os métodos de retreinamento e os critérios de validação que serão aplicados ao longo do ciclo de vida do dispositivo. Dessa forma, um sistema de detecção de retinopatia diabética aprovado inicialmente com dados de uma população específica pode ser legalmente atualizado para melhorar seu desempenho em outras populações, desde que siga o plano de mudanças pré-aprovado, garantindo segurança sem impedir a evolução tecnológica.

4.6 OS PRINCÍPIOS DE IA CONFIÁVEL DA OECD

A governança global de inteligência artificial encontra um dos seus marcos conceituais mais influentes nos Princípios de IA da Organização para a Cooperação e Desenvolvimento Econômico. Adotados em maio de 2019, esses princípios representam o primeiro consenso

multilateral sobre IA entre países, estabelecendo diretrizes éticas e de governança que foram incorporadas por mais de 46 nações e pela UE (OECD, 2019). Embora constituam *soft law*, ou seja, instrumentos não vinculantes juridicamente, os Princípios OECD exercem influência normativa significativa ao moldar legislações nacionais, orientar políticas públicas e fornecer referência para autorregulação corporativa. Para o contexto da saúde, onde sistemas de IA processam dados sensíveis e apoiam decisões clínicas críticas, esses princípios oferecem estrutura conceitual essencial para balancear inovação tecnológica com proteção de direitos fundamentais.

Os cinco princípios estabelecidos pela OECD abordam dimensões complementares da IA confiável. O primeiro princípio, crescimento inclusivo, desenvolvimento sustentável e bem-estar, orienta que sistemas de IA devem beneficiar as pessoas e o planeta ao fomentar crescimento equitativo e sociedades sustentáveis. O segundo, valores centrados no ser humano e equidade, determina que IA deve respeitar direitos fundamentais, valores democráticos e diversidade, protegendo autonomia humana e prevenindo resultados injustos. O terceiro princípio, transparência e explicabilidade, estabelece que deve haver divulgação significativa sobre sistemas de IA para promover compreensão geral e permitir contestação de resultados. O quarto, robustez, segurança e proteção, exige que IA funcione de forma confiável durante todo seu ciclo de vida, com gerenciamento apropriado de riscos. O quinto princípio, responsabilização, determina que organizações e indivíduos desenvolvendo ou operando sistemas de IA devem ser responsabilizáveis por seu funcionamento adequado (OECD, 2019). Para IA em saúde, três desses princípios assumem relevância particular: transparência, que permite aos médicos compreenderem recomendações diagnósticas; robustez, que garante desempenho confiável com dados clínicos imperfeitos; e responsabilização, que define claramente quem responde por erros algoritmos em contextos clínicos.

A relevância desses princípios transcende o plano operacional. Batool et al. (2025) identificaram os Princípios OECD como um dos frameworks de *soft law* mais citados na literatura de governança de IA, funcionando como ponte conceitual entre aspirações éticas abstratas e requisitos técnicos operacionalizáveis. Esta posição os torna exemplo paradigmático de *soft law* que, embora não imponha obrigações jurídicas diretas, exerce poder normativo ao fornecer vocabulário compartilhado e referência para desenvolvimento tanto de *hard law* nacional quanto de padrões técnicos setoriais, estabelecendo conexão direta com a discussão da seção seguinte sobre a interação entre instrumentos vinculantes e orientadores na governança de IA.

4.7 HARD LAW E SOFT LAW NA GOVERNANÇA DE IA

A regulação da inteligência artificial pode ser analisada através da distinção entre *hard law* e *soft law*. Shaffer e Pollack (2010) definiram *hard law* como obrigações juridicamente vinculantes, precisas e com delegação de autoridade para interpretação e aplicação coercitiva, enquanto *soft law* refere-se a arranjos jurídicos enfraquecidos em uma ou mais dessas dimensões (obrigação, precisão e delegação), englobando instrumentos não vinculantes como diretrizes, princípios e recomendações. *Hard law* oferece maior previsibilidade jurídica e mecanismos de aplicação coercitiva, mas impõe custos elevados de negociação e dificuldades de adaptação. *Soft law*, por sua vez, facilita a experimentação e aprendizado em contextos de incerteza tecnológica, permitindo maior flexibilidade sem os custos de revisão de normas vinculantes.

Esses instrumentos podem interagir de formas distintas. Shaffer e Pollack (2010) demonstraram que *hard law* e *soft law* podem operar como complementos, quando *soft law* prepara o caminho para *hard law* ou quando *hard law* é posteriormente elaborado por instrumentos de *soft law*, ou como antagonistas, quando atores promovem *soft law* para minar *hard law* existente, ou vice-versa, especialmente em contextos de conflito distributivo entre Estados e fragmentação de regimes regulatórios.

No campo da governança de IA em saúde, essa distinção se materializa de formas claras. O Regulamento da UE sobre Inteligência Artificial (EU AI Act) classifica sistemas de IA utilizados em dispositivos médicos como de alto risco (Artigo 6, Anexo III), sujeitando-os a requisitos obrigatórios de avaliação de conformidade, documentação técnica, transparência, supervisão humana e robustez (EUROPEAN UNION, 2024). Em contrapartida, abordagens de *soft law* manifestam-se através de princípios e diretrizes não vinculantes, como os Princípios sobre Inteligência Artificial da OECD, que promovem IA centrada no ser humano, transparente e responsável (OECD, 2019), e as diretrizes éticas da Organização Mundial da Saúde para IA em saúde, que orientam o desenvolvimento responsável sem criar obrigações juridicamente exigíveis (WHO, 2021).

Na prática, a interação entre *hard law* e *soft law* na governança de IA frequentemente ocorre de forma complementar. O General Data Protection Regulation (GDPR) estabelece requisitos vinculantes de proteção de dados que se aplicam a sistemas de IA, exigindo que organizações implementem medidas técnicas e organizacionais apropriadas (EUROPEAN UNION, 2016, art. 25). Essa obrigação legal leva organizações a desenvolver códigos de conduta e certificações voluntárias como mecanismos complementares de conformidade

(BATOOL et al., 2025). Tal dinâmica exemplifica como *hard law* pode catalisar o desenvolvimento de *soft law*, criando um modelo de autorregulação monitorada onde instrumentos voluntários preenchem lacunas regulatórias e oferecem flexibilidade para adaptação a inovações tecnológicas. Os conceitos aqui discutidos formam a base analítica para a investigação que se segue dos *frameworks*.

Considerando que a governança de IA em saúde requer articulação entre múltiplos elementos regulatórios, incluindo mecanismos de transparência, estratégias de mitigação de vieses, definição de responsabilidades e escolhas entre instrumentos vinculantes e orientadores, torna-se fundamental examinar como diferentes jurisdições estruturam seus marcos regulatórios. O Brasil, em processo de discussão do PL 2.338/2023, encontra-se em momento oportuno para analisar modelos internacionais consolidados. Nesse contexto, uma análise documental comparativa dos frameworks regulatórios de EUA, EU e Singapura pode identificar abordagens distintas de governança e suas implicações para o desenvolvimento de políticas nacionais adequadas ao contexto brasileiro. A metodologia apresentada a seguir detalha os procedimentos adotados para conduzir essa análise.

5. METODOLOGIA

Este capítulo descreve os procedimentos metodológicos utilizados para atingir os objetivos propostos. A pesquisa adotou abordagem qualitativa, de natureza exploratório-descritiva, com delineamento baseado em estudo documental comparativo.

5.1 DELINEAMENTO DA PESQUISA

Este trabalho foi conduzido como pesquisa documental, método que se baseia em materiais que ainda não receberam tratamento analítico (GIL, 2002). A pesquisa documental se diferencia da bibliográfica por utilizar fontes primárias (documentos "em primeira mão", como leis e relatórios) que ainda não foram submetidas a análise aprofundada, servindo como "fonte natural de informação" (LÜDKE; ANDRÉ, 1986).

Conforme Cellard (2008), a análise documental é adequada para observar a evolução de conceitos e práticas, alinhando-se ao objetivo deste trabalho de compreender o desenvolvimento dos cenários regulatórios. A abordagem foi qualitativa, focando na interpretação e análise crítica do conteúdo dos documentos, não na frequência de ocorrências.

5.2 SELEÇÃO DOS CASOS PARA ANÁLISE COMPARATIVA

A seleção dos casos seguiu abordagem intencional, não probabilística, conhecida como amostragem teórica. Nesta abordagem, os casos não são escolhidos por representatividade estatística, mas por seu potencial de esclarecer diferentes faces do fenômeno e contribuir para compreensão mais abrangente (EISENHARDT, 1989). Os critérios de seleção foram: influência global no debate sobre regulação de tecnologia; maturidade e acessibilidade do arcabouço documental; e representação de abordagens filosóficas distintas à governança de IA.

Com base nesses critérios, foram selecionados quatro casos que funcionam como arquétipos no cenário regulatório, permitindo análise em contrastes e lições aplicáveis ao contexto brasileiro. A justificativa para cada escolha é detalhada no Quadro 1.

Quadro 1 - Cenário regulatório

Modelo de referência	Arquétipo representado	Justificativa de inclusão
Brasil (Anvisa/LGPD/PL) 2.338/2023	Caso em construção	Caso central de análise. País com sistema de saúde universal, em trajetória de convergência com modelo europeu garantista, buscando equilibrar proteção de direitos fundamentais e promoção da inovação em contexto de desigualdades estruturais.
Estados Unidos (FDA)	Inovação adaptativa	Analisar o modelo pioneiro (PCCP) para regulação da evolução de algoritmos, focado no ciclo de vida do produto.
União Europeia (AI Act/GDPR)	Regulação baseada em direitos	Examinar o paradigma horizontal e baseado em risco que influenciou diretamente a proposta legislativa brasileira.
Singapura (IMDA/PDPC)	Governança pragmática	Incluído por representar o arquétipo da governança ágil e orientada ao ecossistema. Diferentemente dos modelos focados primariamente em legislação prescritiva (UE) ou regulação de produto (EUA), Singapura adota uma abordagem de co-regulação, onde o Estado atua como um facilitador, criando ferramentas e frameworks voluntários.

Fonte: Elaborado pelo autor (2025).

5.3 FONTES DE DADOS E DELIMITAÇÃO DO CORPUS

A coleta de dados concentrou-se exclusivamente em fontes primárias oficiais: textos legislativos (leis e regulamentos), normas técnicas (resoluções de agências reguladoras), documentos estratégicos governamentais e diretrizes oficiais emitidas pelas autoridades competentes de cada jurisdição analisada.

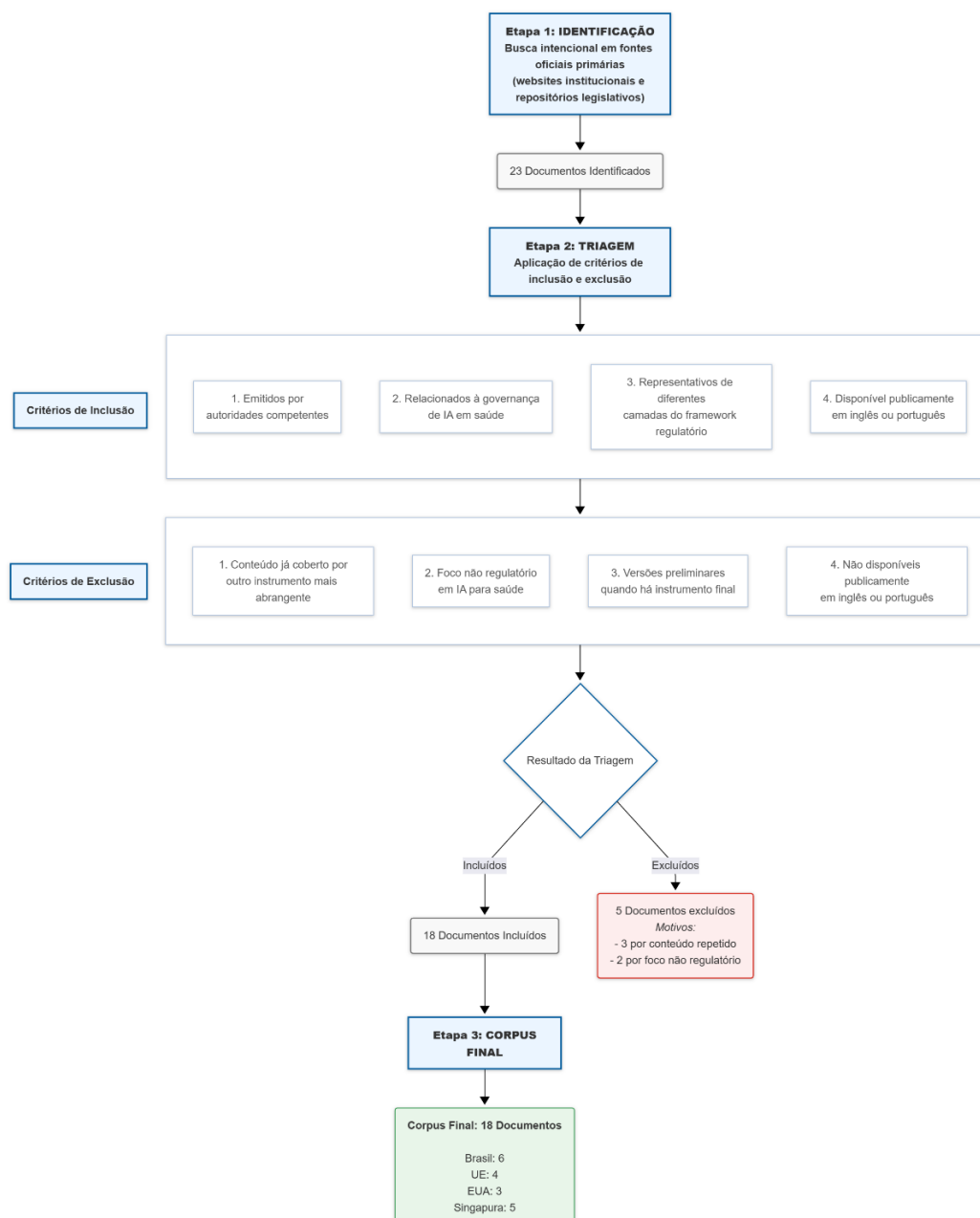
O processo de composição do corpus seguiu três etapas sequenciais. Na etapa de identificação, foi realizada busca intencional nos websites institucionais das autoridades competentes e nos repositórios oficiais de legislação de cada jurisdição selecionada. Para cada país ou região, foram consultados os portais de parlamentos nacionais, agências reguladoras setoriais e órgãos executivos responsáveis por políticas de IA ou saúde digital. Além desses documentos primários, os websites oficiais das respectivas agências reguladoras foram consultados como fontes de apoio para garantir acesso às versões mais recentes e autênticas dos instrumentos normativos. A busca concentrou-se em documentos com potencial pertinência para a governança de IA em saúde. Este procedimento resultou na identificação inicial de 23 documentos oficiais.

Na etapa de triagem, os 23 documentos identificados foram submetidos a leitura preliminar para verificação de elegibilidade segundo critérios de inclusão e exclusão, conforme detalhado na Figura 1. Os critérios de inclusão exigiram que os documentos fossem emitidos por autoridades competentes, relacionados à governança de IA em saúde ou à proteção de dados aplicável a sistemas de IA, representativos de diferentes camadas do *framework* regulatório e publicamente acessíveis em língua inglesa ou portuguesa. Foram excluídos documentos que apresentassem conteúdo já coberto por outro instrumento mais abrangente, foco não regulatório em IA para saúde, caráter preliminar com versão final já publicada ou indisponibilidade em língua acessível. Cabe destacar que o critério de exclusão referente a versões preliminares aplica-se a documentos substituídos por instrumentos finais, como *discussion papers* que deram origem a guias definitivos. Projetos de lei em tramitação avançada, como o PL 2338/2023, foram incluídos por representarem o estado atual do debate regulatório, não havendo instrumento final que os substitua.

A aplicação sistemática desses critérios resultou na exclusão de cinco documentos: três foram excluídos por apresentarem conteúdo já coberto por outros instrumentos mais abrangentes da mesma jurisdição e dois por não abordarem aspectos regulatórios de IA em saúde. O corpus final foi composto por 18 documentos oficiais, distribuídos entre as quatro jurisdições conforme apresentado na Figura 1. Este conjunto documental foi considerado suficiente e representativo para os propósitos da análise comparativa, abrangendo múltiplas

camadas da arquitetura regulatória: estratégica, transversal de proteção de dados, setorial de dispositivos médicos e específica de IA.

Figura 1 - Processo de identificação, triagem e composição do corpus documental



Fonte: Elaborado pelo autor (2025).

5.4 PROCEDIMENTOS DE ANÁLISE

A análise dos dados foi conduzida por meio do método de Análise de Conteúdo Categorical (BARDIN, 2011), desdobrando-se em três fases sequenciais e interdependentes:

Pré-análise: Esta fase compreendeu a organização do material e definição das regras de análise. O protocolo de leitura foi estruturado em torno de cinco categorias analíticas, definidas a priori com base nos Princípios de IA Confiável da OECD (2019):

1. Transparência e explicabilidade;
2. Robustez, segurança e proteção;
3. Responsabilidade (*accountability*);
4. Justiça e privacidade;
5. Governança do ciclo de vida do algoritmo.

Exploração do material: Esta fase consistiu na leitura sistemática e codificação do corpus. Para cada documento, foram extraídas informações pertinentes a cada uma das cinco categorias. O produto desta fase foi organizado como análise descritiva textual (Capítulo 6), na qual a abordagem de cada jurisdição é detalhada para cada categoria, e cada instrumento é classificado quanto à sua natureza jurídica como *hard law* (vinculativo) ou *soft law* (orientador).

Tratamento dos resultados, inferência e interpretação: A fase final concentrou-se na análise comparativa e formulação de conclusões. Primeiramente, os resultados da análise descritiva foram sintetizados em matriz de análise comparativa (Seção 7.1), permitindo visão panorâmica dos diferentes frameworks. Em seguida, foi realizada análise comparativa horizontal dos dados da matriz (Seção 7.2) para identificar convergências e divergências. Por fim, os *insights* derivados desta análise foram interpretados e traduzidos em lições estratégicas e recomendações para a gestão em saúde no Brasil (Capítulo 8), respondendo aos objetivos da pesquisa.

6. RESULTADOS

Conforme detalhado na Figura 1 (seção 5.3), o processo de seleção documental identificou inicialmente 23 documentos oficiais, dos quais cinco foram excluídos após aplicação dos critérios estabelecidos. O corpus final de 18 documentos distribuiu-se entre as quatro jurisdições analisadas: seis do Brasil, quatro da UE, três dos EUA e cinco de Singapura. Os resultados da análise deste corpus são apresentados e discutidos neste capítulo.

A análise descritiva dos *frameworks* regulatórios é exposta na seção 6.1, detalhando a arquitetura de governança e a abordagem de cada jurisdição em relação às categorias analíticas. Subsequentemente, no Capítulo 7, é apresentada a análise comparativa, sintetizada em uma matriz, da qual se extraem as principais convergências e divergências. Por fim, no Capítulo 8, são discutidas as lições estratégicas, os desafios identificados e as recomendações para o

aprimoramento do cenário regulatório brasileiro e para a governança de IA nas instituições de saúde.

6.1 ANÁLISE DESCRITIVA DOS FRAMEWORKS REGULATÓRIOS

Nesta seção, são detalhados os resultados da análise documental de cada um dos quatro frameworks regulatórios, seguindo a estrutura de categorias analíticas. O objetivo é apresentar de forma sistemática as regras, diretrizes e princípios que cada jurisdição estabelece para a governança da IA, com foco especial em suas aplicações na área da saúde.

6.1.1 O FRAMEWORK REGULATÓRIO DO BRASIL

6.1.1.1 ARQUITETURA DO FRAMEWORK E NATUREZA DOS INSTRUMENTOS

A arquitetura regulatória brasileira para IA em saúde caracteriza-se por abordagem que combina legislação consolidada, regulamentos técnicos setoriais e políticas estratégicas orientadoras. Esta estrutura em camadas parte de base da proteção de dados (LGPD) e se especializa progressivamente através de regulação setorial (Anvisa) e futura legislação de IA (PL 2.338/2023). O Quadro 2 detalha os instrumentos que compõem este framework.

Quadro 2 - Arquitetura regulatória do Brasil

Categoria	Instrumento	Natureza Jurídica	Função no Framework
<i>Hard law</i> (vinculativo)	LGPD	Lei federal	Camada fundamental; regula todo o tratamento de dados pessoais no país
<i>Hard law</i> (vinculativo)	RDC Anvisa nº 751/2022	Regulamento técnico setorial	Estabelece as regras para a classificação de risco e regularização de todos os dispositivos médicos
<i>Hard law</i> (vinculativo)	RDC Anvisa nº 657/2022	Regulamento técnico	Complementa a RDC 751; foca nos requisitos específicos para SaMD, incluindo aqueles com IA
<i>Soft law</i> (orientador/estratégico)	PL nº 2.338/2023	Projeto de Lei em tramitação (<i>hard law</i> em potencial)	Representa o consenso político para regulação de IA; quando aprovado, criará obrigações vinculativas para sistemas de alto risco
<i>Soft law</i> (orientador/estratégico)	PBIA	Documento de estratégia nacional	Define a visão e as prioridades do Estado para o fomento e desenvolvimento da IA
<i>Soft law</i> (orientador/estratégico)	Estratégia de Saúde Digital para o Brasil	Política setorial do Ministério da Saúde	Orienta a transformação digital no SUS, focando na governança e interoperabilidade de dados

Fonte: Elaborado pelo autor (2025).

6.1.1.2 ANÁLISE CATEGORIAL

O Quadro 3 apresenta a análise categorial do framework brasileiro, mapeando como cada um dos cinco princípios éticos de IA confiável (transparência e explicabilidade, robustez e segurança, responsabilidade, justiça e privacidade, e governança do ciclo de vida) é operacionalizado através dos diferentes instrumentos normativos e estratégicos do país.

Quadro 3 - Análise categorial do Brasil

Princípio Ético	Instrumento Normativo	Natureza	Disposição Legal e Operacionalização
Transparência e explicabilidade	LGPD	<i>Hard law</i>	Art. 20: garante o direito de solicitar revisão de decisões tomadas unicamente com base em tratamento automatizado. § 1º: exige que o controlador forneça informações sobre os "critérios e dos procedimentos utilizados para a decisão"
Transparência e explicabilidade	PL nº 2.338/2023	<i>Soft law</i> (futura <i>Hard law</i>)	Artigo 3º, VI: consagra a "transparência e explicabilidade" como princípio. Artigo 6º, I: cria um "direito à explicação sobre a decisão" para sistemas de alto risco. Artigo 7º: deve ser fornecida em linguagem simples e acessível
Transparência e explicabilidade	RDC Anvisa nº 657/2022	<i>Hard law</i>	Artigo 7º, III: determina que as instruções de uso do software devem incluir "descrições genéricas dos algoritmos, rotinas e fórmulas utilizadas para gerar o processamento clínico"
Transparência e explicabilidade	PBIA	<i>Soft law</i>	Ação 51: propõe a criação de um "Centro Nacional de Transparência Algorítmica", sinalizando uma preocupação estrutural com o tema
Transparência e explicabilidade	Estratégia de Saúde Digital	<i>Soft law</i>	Materializa a transparência ao propor a implementação de um registro pessoal de saúde que garanta ao cidadão o acesso e o controle sobre seus próprios dados
Robustez, segurança e proteção	LGPD	<i>Hard law</i>	Artigo 6º, VII: estabelece o princípio da "segurança". Artigo 46: obrigação de adotar "medidas de segurança, técnicas e administrativas", que devem ser consideradas desde a concepção do produto
Robustez, segurança e proteção	PL nº 2.338/2023	<i>Soft law</i> (futura <i>Hard law</i>)	Artigo 3º, VIII: eleva a "confiabilidade e robustez" a princípio orientador. Artigo 18: exige que os agentes de IA de alto risco realizem "testes para avaliação de níveis apropriados de segurança"
Robustez, segurança e proteção	RDC Anvisa nº 751/2022	<i>Hard law</i>	Centra-se na segurança via gestão de risco e na elaboração de um Dossiê Técnico que comprove o desempenho seguro
Robustez, segurança e proteção	RDC Anvisa nº 657/2022	<i>Hard law</i>	Arts. 2º, III (define cibersegurança); 7º, VI (exige informações nas instruções de uso); 11 (exige informações na arquitetura do produto); e 17, III (determina desenvolvimento considerando gestão de risco, segurança das informações, verificação e validação)
Responsabilidade (accountability)	LGPD	<i>Hard law</i>	Artigo 6º, X: consagra o princípio da "responsabilização e prestação de contas", exigindo que os agentes sejam capazes de demonstrar a adoção de medidas eficazes, sendo o DPIA o principal instrumento. Artigo 42: estabelece regime

Princípio Ético	Instrumento Normativo	Natureza	Disposição Legal e Operacionalização
			de responsabilidade solidária entre controlador e operador
Responsabilidade (accountability)	PL nº 2.338/2023	<i>Soft law</i> (futura <i>Hard law</i>)	Capítulo V: prevê capítulo específico para a Responsabilidade Civil. Artigo 37: inversão do ônus da prova em favor da vítima, tornando os desenvolvedores e aplicadores proativamente responsáveis por provar a segurança de seus sistemas em caso de dano
Responsabilidade (accountability)	RDC Anvisa nº 751/2022	<i>Hard law</i>	Artigo 4º: atribui a responsabilidade primária ao "detentor" do registro e ao "responsável técnico", criando uma dupla camada de accountability corporativa e profissional
Justiça e privacidade	LGPD	<i>Hard law</i>	Artigo 1º: a privacidade é a espinha dorsal da LGPD, qualificando "dados de saúde" como sensíveis, sujeitos a regras mais estritas de tratamento. Artigo 6º, IX: codifica a justiça como o princípio da "não discriminação"
Justiça e privacidade	PL nº 2.338/2023	<i>Soft law</i> (futura <i>Hard law</i>)	Artigo 5º, III: estabelece um "direito à correção de vieses discriminatórios". Artigo 18: obriga agentes de alto risco a adotar "medidas para mitigar e prevenir vieses"
Justiça e privacidade	Estratégia de Saúde Digital e PBIA	<i>Soft law</i>	Posicionam a justiça como objetivo final, visando reduzir desigualdades e garantir que a IA beneficie toda a sociedade
Justiça e privacidade	RDC Anvisa nº 751/2022 e nº 657/2022	<i>Hard law</i>	Abordam o tema de forma implícita, através da exigência de validação clínica que garanta a eficácia do dispositivo para toda a população-alvo
Governança do ciclo de vida do algoritmo	LGPD	<i>Hard law</i>	Artigos 15 e 16: impõe o princípio da "finalidade" e a obrigação de eliminação dos dados após o término do tratamento, limitando indiretamente o ciclo de vida do algoritmo
Governança do ciclo de vida do algoritmo	PL nº 2.338/2023	<i>Soft law</i> (futura <i>Hard law</i>)	Artigo 4º, II: define "ciclo de vida". Artigo 3º, VII: exige "diligência devida e auditabilidade ao longo de todo o ciclo de vida". Artigo 26: a Avaliação de Impacto Algorítmico é concebida como um "processo iterativo contínuo"
Governança do ciclo de vida do algoritmo	RDC Anvisa nº 751/2022	<i>Hard law</i>	Estabelece a governança por meio da revalidação do registro e da necessidade de aprovação de alterações
Governança do ciclo de vida do algoritmo	RDC Anvisa nº 657/2022	<i>Hard law</i>	Artigo 16: estipula que alterações que "afetem significativamente as funcionalidades clínicas, segurança e eficácia" demandam nova análise da Anvisa, criando um controle de mudanças pós-mercado

Fonte: Elaborado pelo autor (2025).

6.1.2 O FRAMEWORK REGULATÓRIO DA UE

6.1.2.1 ARQUITETURA DO FRAMEWORK E NATUREZA DOS INSTRUMENTOS

O framework regulatório da UE é o mais integrado e legalmente prescritivo entre os casos estudados. Sua arquitetura é marcadamente *top-down* e se consolida por meio de

regulamentos com força de lei direta, aplicáveis uniformemente em todos os Estados-Membros, visando à criação de um mercado único digital. Essa abordagem resulta em um sistema robusto e com pouca distinção entre *hard* e *soft law*, onde mesmo os documentos orientadores servem primariamente para harmonizar a aplicação das leis já existentes. O Quadro 4 a seguir detalha os instrumentos que compõem este framework.

Quadro 4 - Arquitetura regulatória da UE

Categoria	Instrumento	Natureza Jurídica	Função no Framework
<i>Hard law</i> (vinculativo)	Regulamento (UE) 2016/679 (GDPR)	Regulamento pan-europeu com força de lei direta	Camada fundamental de proteção de dados; base para os direitos digitais dos cidadãos europeus
<i>Hard law</i> (vinculativo)	Regulamento (UE) 2017/745 (MDR)	Regulamento pan-europeu com força de lei direta	Camada setorial de saúde; requisitos de segurança e ciclo de vida para dispositivos médicos, incluindo softwares (SaMD)
<i>Hard law</i> (vinculativo)	Regulamento (UE) 2024/1689 (AI Act)	Regulamento pan-europeu com força de lei direta	Camada de tecnologia; primeiro marco legal horizontal específico para IA; abordagem baseada em risco com obrigações para sistemas de alto risco
<i>Soft law</i> (orientador/estratégico)	Interplay between the MDR & IVDR and the AIA	Guia orientador	Documento de harmonização interpretativo; explicação da aplicação conjunta do MDR e AI Act; não cria novas obrigações

Fonte: Elaborado pelo autor (2025).

6.1.2.2 ANÁLISE CATEGORIAL

O Quadro 5 apresenta a análise categorial do framework da UE, demonstrando como os princípios éticos de IA confiável são traduzidos em obrigações explícitas e vinculativas, distribuídas de forma complementar entre três regulamentos principais: GDPR (proteção de dados), MDR (dispositivos médicos) e AI Act (sistemas de IA). Esta estrutura integrada busca criar um ecossistema regulatório sem lacunas.

Quadro 5 - Análise categorial da UE

Princípio Ético	Instrumento Normativo	Natureza	Disposição Legal e Operacionalização
Transparência e explicabilidade	GDPR	<i>Hard law</i>	Art. 22: confere o direito de contestar decisões automatizadas e obter intervenção humana, interpretado como a gênese do 'direito à explicação'
Transparência e explicabilidade	AI Act	<i>Hard law</i>	Art. 13: exige que sistemas de alto risco sejam 'suficientemente transparentes' para permitir interpretação adequada dos resultados. Art. 11: impõe documentação técnica detalhada. Art. 12: exige capacidade de manter registros automáticos (<i>logs</i>), criando trilha de auditoria para fins de explicabilidade posterior

Princípio Ético	Instrumento Normativo	Natureza	Disposição Legal e Operacionalização
Transparência e explicabilidade	MDR	<i>Hard law</i>	Anexo I, Seção 23: exige informações detalhadas nas instruções de utilização para garantir transparência ao usuário final
Transparência e explicabilidade	Guia AIB/MDCG	<i>Soft law</i>	Consolida os conceitos, afirmando que os regulamentos tornam transparência e explicabilidade 'obrigações vinculativas'
Robustez, segurança e proteção	MDR	<i>Hard law</i>	Anexo I: define Requisitos Gerais de Segurança e Desempenho, exigindo que <i>softwares</i> garantam 'repetibilidade, fiabilidade e desempenho'. Seção 17.4: aborda cibersegurança, incluindo proteção contra acesso não autorizado
Robustez, segurança e proteção	GDPR	<i>Hard law</i>	Artigo 32: exige a implementação de "medidas técnicas e organizativas adequadas". Artigo 25: promove a "segurança por concepção e por defeito"
Robustez, segurança e proteção	AI Act	<i>Hard law</i>	Art. 15: exige 'exatidão, solidez e cibersegurança', abordando vulnerabilidades específicas da IA como proteção de dados de treinamento e prevenção de envenenamento de modelos
Robustez, segurança e proteção	Guia AIB/MDCG	<i>Soft law</i>	Destaca a integração da robustez do dispositivo, a segurança dos dados e a cibersegurança como obrigações legais interligadas
Responsabilidade (accountability)	GDPR	<i>Hard law</i>	Art. 5(2): estabelece o 'princípio da responsabilidade', invertendo o ônus da prova: o controlador deve demonstrar cumprimento da lei. Operacionalizado pela obrigação de realizar DPIA
Responsabilidade (accountability)	AI Act	<i>Hard law</i>	Distribui obrigações ao longo da cadeia entre 'prestadores' (desenvolvedores, Art. 16) e 'responsáveis pela implantação' (usuários, Art. 26). Art. 12: garante demonstração de responsabilidade via documentação técnica e rastreabilidade (<i>logs</i>)
Responsabilidade (accountability)	MDR	<i>Hard law</i>	Artigo 15: cria a figura da "Pessoa responsável pela observância da regulamentação", estabelecendo um ponto focal de accountability dentro de cada fabricante
Responsabilidade (accountability)	Guia AIB/MDCG	<i>Soft law</i>	Aponta a documentação técnica e rastreabilidade via logs como mecanismos de demonstração de responsabilidade
Justiça e privacidade	GDPR	<i>Hard law</i>	A privacidade é direito fundamental e razão de ser do GDPR. Art. 9: classifica 'dados relativos à saúde' como categoria especial, cujo tratamento é proibido por padrão salvo sob condições estritas
Justiça e privacidade	AI Act	<i>Hard law</i>	Art. 10 ('Dados e governação de dados'): transforma a mitigação de viés em obrigação legal. Exige que conjuntos de dados de treinamento sejam 'pertinentes, suficientemente representativos e isentos de erros'. Exige exame ativo para detecção de vieses e adoção de medidas de mitigação
Justiça e privacidade	Guia AIB/MDCG	<i>Soft law</i>	Torna a monitorização de vieses uma exigência legal para garantir que os sistemas de IA não levem a discriminações
Governança do ciclo de vida do algoritmo	MDR	<i>Hard law</i>	Anexo I: exige Sistema de Gestão de Risco como 'processo iterativo contínuo'. Art. 83: exige Sistema de Monitorização Pós-Comercialização para supervisão do produto após lançamento
Governança do ciclo de vida do algoritmo	GDPR	<i>Hard law</i>	Impõe a "limitação da finalidade" e a "limitação da conservação", impactando diretamente a longevidade dos algoritmos

Princípio Ético	Instrumento Normativo	Natureza	Disposição Legal e Operacionalização
Governança do ciclo de vida do algoritmo	AI Act	Hard law	Art. 9: exige Sistema de Gestão de Riscos abrangendo 'todo o ciclo de vida do sistema de IA'. Art. 72: exige sistema de acompanhamento pós-comercialização. Art. 3, item 23: 'modificação substancial' funciona como gatilho regulatório que força nova avaliação de conformidade, garantindo supervisão contínua

Fonte: Elaborado pelo autor (2025).

6.1.3 O FRAMEWORK REGULATÓRIO DO EUA

6.1.3.1 ARQUITETURA DO FRAMEWORK E NATUREZA DOS INSTRUMENTOS

O framework regulatório dos EUA distingue-se por abordagem descentralizada, pró-inovação e predominantemente orientada por *soft law*. Em vez de lei federal abrangente, a governança é construída a partir de diretrizes estratégicas do Poder Executivo (Executive Order), frameworks técnicos voluntários (NIST AI RMF) e guias setoriais de agências reguladoras (FDA). Esta arquitetura flexível busca adaptar-se ao avanço tecnológico. O Quadro 6 detalha os principais instrumentos.

Quadro 6 - Arquitetura regulatória do EUA

Categoria	Instrumento	Natureza Jurídica	Função no Framework
<i>Soft law</i> (orientador/estratégico)	Executive Order 14110 ("Safe, Secure, and Trustworthy AI")	Diretiva presidencial	Funciona como o "guarda-chuva" estratégico; estabelece a política nacional e ordena que as agências federais desenvolvam padrões e guias, mas não cria obrigações diretas para o setor privado
<i>Soft law</i> (orientador/estratégico)	NIST AI Risk Management Framework (AI RMF 1.0)	Framework voluntário	De altíssima influência; fornece a metodologia padrão do governo e da indústria para a gestão de riscos de IA; principal guia de governança
<i>Soft law</i> (orientador/estratégico)	Guias da FDA (Discussion Paper e PCCP Guidance)	Documentos de orientação	Descrevem como a FDA pretende regular os dispositivos médicos com IA. Embora não sejam legalmente vinculativos, representam a interpretação da agência e sinalizam as expectativas para a aprovação de produtos. A adesão a eles é, na prática, o caminho para a conformidade

Fonte: Elaborado pelo autor (2025).

6.1.3.2 ANÁLISE CATEGORIAL

O Quadro 7 apresenta a análise categorial do framework dos EUA, demonstrando como a governança de IA é construída através da articulação de instrumentos de *soft law*. O quadro

revela uma cascata de influência: prioridades políticas (Executive Order) são convertidas em metodologia de gestão de risco (NIST AI RMF), aplicada setorialmente nos guias da FDA. Esta estrutura enfatiza gestão de riscos, boas práticas de engenharia e governança interna.

Quadro 7 - Análise categorial do EUA

Princípio Ético	Instrumento Normativo	Natureza	Disposição Legal e Operacionalização
Transparência e explicabilidade	NIST AI RMF	<i>Soft law</i>	Define 'Explicabilidade e Interpretabilidade' como característica da IA Confiável. Defende nível de explicabilidade 'adequado ao propósito', gerenciado com base no risco e público-alvo, não como requisito absoluto
Transparência e explicabilidade	FDA (PCCP Guidance)	<i>Soft law</i>	Exige atualização contínua da rotulagem do dispositivo para informar usuários sobre modificações implementadas e seus impactos no desempenho
Transparência e explicabilidade	Executive Order 14110	<i>Soft law</i>	Seção 4.5: prioriza criação de mecanismos de identificação de conteúdo sintético (como marca d'água) para combater desinformação
Robustez, segurança e proteção	Executive Order 14110	<i>Soft law</i>	Seção 2a: estabelece que 'IA deve ser segura e protegida', delegando ao NIST estabelecer diretrizes e padrões para testes robustos, incluindo metodologia de teste cibernético com hackers éticos. Seção 4.1: visa identificar vulnerabilidades sistêmicas para aprimoramento de segurança
Robustez, segurança e proteção	NIST AI RMF	<i>Soft law</i>	Trata 'Seguro, Protegido e Resiliente' como características gerenciáveis de IA Confiável, fornecendo metodologia para identificação e mitigação de riscos
Robustez, segurança e proteção	FDA (PCCP Guidance)	<i>Soft law</i>	Garante a segurança futura dos algoritmos de aprendizado contínuo ao focar na robustez do "Protocolo de Modificação", que deve conter critérios de aceitação pré-definidos e métodos de validação rigorosos para cada mudança, garantindo que o dispositivo "permanecerá seguro e eficaz"
Responsabilidade (accountability)	NIST AI RMF	<i>Soft law</i>	Dedica sua função "Governar" inteiramente a este princípio, enfatizando a importância de definir "funções e responsabilidades" claras e de estabelecer políticas internas para a supervisão da gestão de riscos. A responsabilidade é vista como uma consequência de uma boa governança interna
Responsabilidade (accountability)	Executive Order 14110	<i>Soft law</i>	Seção 10: determina que cada agência federal designe um <i>Chief AI Officer</i> responsável por definir, liderar e supervisionar a estratégia de IA
Responsabilidade (accountability)	FDA (PCCP Guidance e Discussion Paper)	<i>Soft law</i>	A responsabilidade do fabricante concentra-se na fase pré-mercado. O fabricante tem ônus de provar segurança e eficácia de seu produto e PCCP para obter autorização, tornando-se responsável por seguir o plano aprovado
Justiça e privacidade	Executive Order 14110	<i>Soft law</i>	Seção 2d: declara que administração 'não irá tolerar uso de IA para desfavorecer' grupos vulneráveis. Seção 8.b: criar uma força-tarefa do HHS para promover 'uso equitativo de IA na saúde'. Seção 9: aborda privacidade como risco a mitigar via PETs

Princípio Ético	Instrumento Normativo	Natureza	Disposição Legal e Operacionalização
Justiça e privacidade	NIST AI RMF	<i>Soft law</i>	Traduz prioridade política de justiça em metodologia técnica, tratando 'gestão de viés' como risco técnico a mapear, medir e gerir
Justiça e privacidade	FDA (PCCP Guidance)	<i>Soft law</i>	Exige que as práticas de gestão de dados do protocolo abordem explicitamente como obter dados representativos de diversas populações e quais as "estratégias para entender e mitigar vieses"
Governança do ciclo de vida do algoritmo	FDA (PCCP Guidance e Discussion Paper de 2019)	<i>Soft law</i>	O PCCP (<i>Predetermined Change Control Plan</i>) é plano de governança do ciclo de vida. Permite que algoritmos de aprendizado contínuo evoluam no mercado dentro de 'envelope' pré-aprovado, desde que fabricante siga rigorosamente 'Protocolo de Modificação', representando solução para regular tecnologias dinâmicas
Governança do ciclo de vida do algoritmo	NIST AI RMF	<i>Soft law</i>	Fornece metodologia de gestão de riscos contínua e interativa, aplicada 'ao longo de todo o 'ciclo de vida da IA'

Fonte: Elaborado pelo autor (2025).

6.1.4 O FRAMEWORK REGULATÓRIO DE SINGAPURA

6.1.4.1 ARQUITETURA DO FRAMEWORK E NATUREZA DOS INSTRUMENTOS

O framework regulatório de Singapura apresenta abordagem pragmática e orientada ao ecossistema, onde a governança é concebida como vantagem competitiva para fomentar inovação. Sua arquitetura híbrida combina leis de base (*hard law*) para proteção de dados e segurança de produtos com guias e estratégias de *soft law* para temas emergentes de IA. Esta estrutura equilibra segurança jurídica e agilidade. O Quadro 8 detalha os instrumentos que formam este framework.

Quadro 8 - Arquitetura regulatória de Singapura

Categoria	Instrumento	Natureza Jurídica	Função no Framework
<i>Hard law</i> (vinculativo)	Health Products (Medical Devices) Regulations 2010	Legislação	Camada de regulação de saúde; foca na segurança do produto e na vigilância de mercado
<i>Hard law</i> (vinculativo)	Personal Data Protection Act 2012 (PDPA)	Lei	Camada de regulação de dados; estabelece as obrigações gerais para o tratamento de dados pessoais em todos os setores
<i>Soft law</i> (orientador/estratégico)	NAIS 2.0 (National AI Strategy)	Documento de estratégia nacional	Define a visão nacional de IA e enquadra governança de confiança como pilar para sucesso econômico; define a visão de Singapura e enquadra a governança de confiança como um pilar para o sucesso econômico
<i>Soft law</i> (orientador/estratégico)	Model AI Governance Framework	Guia voluntário e transversal	Introduz conceitos inovadores como a "responsabilidade compartilhada" e os

			"rótulos informativos"; principal ferramenta de governança de IA do país
Soft law (orientador/estratégico)	Artificial Intelligence in Healthcare Guidelines (AIHGle)	Guia de melhores práticas	Guia específico para o setor da saúde; detalha a implementação da IA em ambientes clínicos de forma ágil e segura

Fonte: Elaborado pelo autor (2025).

6.1.4.2 ANÁLISE CATEGORIAL

O Quadro 9 apresenta a análise categorial do framework de Singapura, ilustrando sua abordagem pragmática. O quadro demonstra como o país utiliza *soft law* (Model AI Governance Framework e AIHGle) para introduzir conceitos como 'responsabilidade compartilhada' e 'rótulos informativos', sustentados por base de *hard law* (PDPA e Health Products Regulations) que garante direitos e deveres fundamentais, criando modelo de co-regulação.

Quadro 9 - Análise categorial de Singapura

Princípio Ético	Instrumento Normativo	Natureza	Disposição Legal e Operacionalização
Transparência e explicabilidade	Model AI Governance Framework	<i>Soft law</i>	Propõe transparência através de 'rótulos informativos': modelos de IA acompanhados de documentação clara sobre dados utilizados, desempenho em testes e limitações, capacitando atores da cadeia a tomar decisões informadas
Transparência e explicabilidade	AIHGle	<i>Soft law</i>	Distingue transparência (informar ao usuário final que interage com IA) de explicabilidade (necessária para decisões clínicas, contextual e adequada ao público: médico vs. paciente), recomendando técnicas de interpretabilidade para modelos complexos
Transparência e explicabilidade	PDPA	<i>Hard law</i>	Fornecer base legal para transparência através da 'Obrigação de Notificação' (informar propósitos do tratamento) e 'Obrigação de Acesso' (saber como dados foram usados)
Robustez, segurança e proteção	Model AI Governance Framework	<i>Soft law</i>	Incentiva testagem e auditoria por terceiros como mecanismo de mercado para criar confiança, recomendando testes cibernéticos com hackers éticos para avaliar robustez. Adapta o princípio de segurança por concepção, destacando necessidade de mecanismos específicos para IA, como filtros de prompts maliciosos
Robustez, segurança e proteção	AIHGle	<i>Soft law</i>	Aborda cibersegurança e proteção de dados, exigindo que software seja capaz de 'prevenir, detectar, responder e recuperar-se de riscos', alinhando-se com diretrizes da autoridade de proteção de dados
Robustez, segurança e proteção	Health Products Regulations	<i>Hard law</i>	Parte VIII: garante segurança via processo reativo de monitoramento e reporte obrigatório de defeitos e eventos adversos
Responsabilidade (accountability)	Model AI Governance Framework	<i>Soft law</i>	Articula que responsabilidade deve ser alocada ao longo da cadeia de valor com base no 'nível de controle' de cada ator

Princípio Ético	Instrumento Normativo	Natureza	Disposição Legal e Operacionalização
Responsabilidade (accountability)	AIHGle	<i>Soft law</i>	Recomenda uso de SLAs para formalizar contratualmente divisão de responsabilidades entre desenvolvedor e implementador da IA
Responsabilidade (accountability)	PDPA	<i>Hard law</i>	Estabelece 'Obrigação de Responsabilidade', exigindo implementação de políticas internas e designação de <i>Data Protection Officer</i> (DPO)
Responsabilidade (accountability)	Health Products Regulations	<i>Hard law</i>	Impõem deveres legais explícitos a fabricantes e importadores
Justiça e privacidade	PDPA	<i>Hard law</i>	Funciona de forma similar ao GDPR, com obrigações de consentimento, limitação de propósito e proteção de dados
Justiça e privacidade	Model AI Governance Framework	<i>Soft law</i>	Aborda justiça como questão de higiene e qualidade de dados. Recomenda adoção de 'medidas de controle de qualidade' e uso de ferramentas para 'remoção de vies
Justiça e privacidade	AIHGle	<i>Soft law</i>	Transforma em requisito de engenharia, exigindo na fase de design garantia de 'representatividade dos <i>datasets</i> ' e teste de desempenho em diferentes subgrupos demográficos para mitigar vieses
Governança do ciclo de vida do algoritmo	AIHGle	<i>Soft law</i>	Estrutura orientação nas fases de Desenvolvimento (<i>Design, Build, Test</i>) e Implementação (<i>Use, Monitor, Review</i>). Aborda algoritmos de aprendizado contínuo recomendando implantação inicial de modelos 'bloqueados' que aprendem 'em paralelo' antes de atualizações, abordagem semelhante ao PCCP da FDA
Governança do ciclo de vida do algoritmo	Model AI Governance Framework	<i>Soft law</i>	Contribui com governança pós-implementação destacando importância de sistemas de reporte de incidentes, criando feedback loop para melhoria contínua
Governança do ciclo de vida do algoritmo	Health Products Regulations e PDPA	<i>Hard law</i>	Governam ciclo de vida reativamente, através de vigilância pós-mercado e 'obrigação de limitação de retenção' de dados

Fonte: Elaborado pelo autor (2025).

7. ANÁLISE COMPARATIVA

Após a caracterização individual dos quatro frameworks regulatórios no capítulo anterior, este capítulo apresenta a análise comparativa sistemática dos dados coletados. O objetivo é identificar padrões transnacionais, divergências estratégicas e as diferentes filosofias de governança que emergem do contraste entre os modelos. A análise está estruturada em três etapas: primeiro, a matriz comparativa sintetiza os achados de forma panorâmica; em seguida, a discussão categorial aprofunda as convergências e divergências para cada uma das cinco dimensões analisadas; por fim, a síntese integra os achados e identifica os arquétipos regulatórios principais. A discussão das implicações estratégicas desses achados para o contexto brasileiro será desenvolvida no capítulo seguinte.

7.1 MATRIZ DE ANÁLISE COMPARATIVA

Quadro 10 – Análise comparativa dos instrumentos

Jurisdição	Síntese da abordagem	Natureza predominante (Hard/Soft law)
Transparência e Explicabilidade		
União Europeia (UE)	Abordagem garantista e prescritiva, baseada em direitos legalmente exigíveis. O cidadão tem o direito de contestar decisões automatizadas e obter uma explicação da lógica envolvida (GDPR, Art. 22). Sistemas de IA de alto risco são legalmente obrigados a ser "suficientemente transparentes" e a possuir capacidade técnica de registrar eventos automaticamente (logs) para auditoria e rastreabilidade (AI Act, Arts. 12-13). A transparência é uma obrigação legal vinculativa.	<i>Hard law.</i>
Estados Unidos (EUA)	Abordagem orientada a riscos, contextual e focada na governança interna e na comunicação com o mercado. A explicabilidade deve ser "adequada ao propósito" e ao risco (NIST AI RMF). No setor de saúde, o foco é na transparência evolutiva, obrigando fabricantes a informar usuários sobre modificações em algoritmos (FDA PCCP Guidance). Há ênfase política na identificação da proveniência de conteúdo gerado por IA (Executive Order).	<i>Soft law.</i>
Singapura	Abordagem pragmática e informativa, focada na capacitação da cadeia de valor. Recomenda "rótulos informativos" sobre os modelos de IA, detalhando dados de treino e limitações (Model AI Governance Framework). Distingue entre a transparência para o paciente (saber que usa IA) e a explicabilidade para o clínico (entender a lógica), sendo a última contextual (AIHGle).	Mista (predominância de <i>Soft law</i>).
Brasil	Abordagem garantista, baseada em direitos já estabelecidos e em consolidação. A LGPD (Art. 20) garante o direito à revisão de decisões automatizadas. A RDC 657/2022 inova ao exigir a descrição de "algoritmos e fórmulas" nas instruções de uso. O futuro Marco Legal (PL 2.338/2023) propõe elevar a explicação a um direito explícito.	Mista (forte tendência a <i>Hard law</i>).
Robustez, segurança e proteção		
União Europeia (UE)	Abordagem integrada e prescritiva, com obrigações legais em todas as camadas. A segurança do produto é o foco do MDR, a segurança dos dados é garantida pelo GDPR, e a robustez da IA é uma exigência explícita do AI Act (Art. 15).	<i>Hard law.</i>
Estados Unidos (EUA)	Abordagem orientada a processos, focada em boas práticas de engenharia. A robustez é alcançada pela adesão a processos de alta qualidade, como a metodologia de gestão de riscos do NIST AI RMF e, na saúde, pela robustez do "Protocolo de Modificação" definido no PCCP da FDA.	<i>Soft law.</i>
Singapura	Abordagem pragmática e verificável, focada na confiança do mercado. A segurança é comprovada via testagem e auditoria por terceiros (<i>soft law</i>) e garantida por um sistema reativo de reporte obrigatório de incidentes (<i>hard law</i>).	Mista (predominância de <i>Soft law</i>).
Brasil	Abordagem fortemente influenciada por padrões internacionais. A regulação da Anvisa (RDC 657) é o pilar central, exigindo a demonstração de conformidade com padrões de engenharia de software e gestão de risco (IEC 62304, ISO 14971) através de um Dossiê Técnico.	Mista (forte tendência a <i>Hard Law</i>).

Jurisdição	Síntese da abordagem	Natureza predominante (Hard/Soft law)
Responsabilidade (accountability)		
União Europeia (UE)	Abordagem jurídica, proativa e de responsabilização em cadeia. Exige a capacidade de demonstrar conformidade (GDPR), distribui obrigações legais entre "prestadores" e "usuários" (AI Act) e personaliza a accountability na figura da "Pessoa Responsável" (MDR).	<i>Hard law.</i>
Estados Unidos (EUA)	Abordagem organizacional, centrada na governança interna e na responsabilidade pré-mercado. O foco é criar uma "cultura de risco" (NIST AI RMF) e responsabilizar o fabricante no processo de aprovação, onde o ônus da prova é dele (FDA).	<i>Soft law.</i>
Singapura	Abordagem inovadora, baseada em responsabilidade compartilhada e contratual. A accountability é alocada com base no "nível de controle" de cada ator e formalizada por meio de contratos (SLAs) entre desenvolvedores e implementadores.	Mista (predominância de <i>Soft law</i>).
Brasil	Abordagem garantista e juridicamente robusta, com forte foco na responsabilização por danos. Além da responsabilidade solidária (LGPD) e da figura do Responsável Técnico (Anvisa), destaca-se pela proposta de inversão do ônus da prova em favor da vítima (PL 2.338/2023).	Mista (forte tendência a <i>Hard law</i>).
Justiça e privacidade		
União Europeia (UE)	Abordagem de vanguarda, com obrigações legais explícitas. A Privacidade é garantida pelo GDPR. A Justiça é transformada em uma obrigação legal de engenharia pelo AI Act (Art. 10), que exige dados de treino representativos e a mitigação ativa de vieses.	<i>Hard law.</i>
Estados Unidos (EUA)	Abordagem de direitos civis e gestão de riscos. A Justiça é uma prioridade política (Executive Order) traduzida em uma metodologia de gestão de risco técnico (NIST, FDA). Carece de uma lei federal de Privacidade abrangente e de <i>hard law</i> específica para viés algorítmico.	<i>Soft law.</i>
Singapura	Abordagem técnica, tratando justiça como qualidade de dados. A Privacidade é garantida pela PDPA (<i>hard law</i>). A Justiça é abordada como uma boa prática de engenharia, recomendando a "remoção de viés (debiasing)" nos dados (<i>soft law</i>).	Mista.
Brasil	Abordagem robusta, alicerçada em uma forte lei de dados. A Privacidade é garantida pela LGPD (<i>hard law</i>), que também estabelece a "não discriminação" como princípio. O futuro Marco Legal visa a criar um direito explícito à correção de vieses.	Mista (forte tendência a <i>Hard law</i>).
Governança do ciclo de vida do algoritmo		
União Europeia (UE)	Abordagem de supervisão contínua e legalmente vinculativa. Exige gestão de risco e monitoramento pós-mercado durante todo o ciclo de vida (MDR, AI Act). O conceito de "modificação substancial" funciona como um gatilho legal para forçar uma nova avaliação de conformidade.	<i>Hard law.</i>
Estados Unidos (EUA)	Abordagem de regulação adaptativa e prospectiva. A principal inovação é o PCCP da FDA, que aprova a priori um protocolo de evolução para o algoritmo, permitindo mudanças controladas sem nova submissão. É uma governança do processo de mudança, não do produto estático.	<i>Soft law.</i>

Jurisdição	Síntese da abordagem	Natureza predominante (Hard/Soft law)
Singapura	Abordagem ágil, baseada em boas práticas de engenharia de software. Orienta a governança nas fases de Desenvolvimento e Implementação (AIHGle), recomendando "aprendizado em paralelo" e reporte de incidentes para melhoria contínua (Model Framework).	Mista (predominância de <i>Soft law</i>).
Brasil	Abordagem mista, combinando controle reativo com avaliação contínua. A Anvisa (RDC 657) já exige nova análise para alterações significativas (<i>hard law</i>). O PL 2.338/2023 propõe uma Avaliação de Impacto Algorítmico contínua como obrigação legal (futura <i>hard law</i>).	Mista (forte tendência a <i>Hard law</i>).

Fonte: Elaborado pelo autor (2025).

A classificação 'Mista' indica coexistência de instrumentos *hard law* e *soft law*, com a indicação de 'predominância' ou 'forte tendência' sinalizando o vetor de desenvolvimento regulatório da jurisdição. A análise baseia-se nos instrumentos normativos identificados nos Quadros 1-8 e considera tanto os instrumentos vigentes quanto aqueles em fase avançada de tramitação (como o PL 2.338/2023 no Brasil), que representam a direção futura da regulação. A matriz evidencia diferentes filosofias regulatórias: a UE privilegia obrigações vinculativas (*hard law*); os EUA adotam abordagem orientada por padrões técnicos e governança voluntária (*soft law*); Singapura e Brasil apresentam arquiteturas mistas em evolução. As categorias de análise seguem o framework conceitual dos princípios éticos de IA confiável, com análise detalhada de cada categoria apresentada nas seções 6.1.1.2 a 6.1.4.2.

7.2 DISCUSSÃO DAS CONVERGÊNCIAS E DIVERGÊNCIAS ESTRATÉGICAS

A análise comparativa dos quatro frameworks regulatórios, sintetizada no Quadro 9, permite identificar padrões de convergência e divergência que transcendem as particularidades nacionais, revelando filosofias regulatórias distintas sobre o papel da IA na sociedade. Esta seção explora essas diferenças estratégicas, categoria por categoria, extraíndo arquétipos que podem orientar tanto o desenvolvimento de políticas públicas quanto a compreensão da governança global da IA em saúde. A estrutura de cada subseção segue o mesmo padrão analítico: identificar a convergência, aprofundar a divergência estratégica e apontar os arquétipos que emergem da comparação.

7.2.1 TRANSPARÊNCIA E EXPLICABILIDADE

A análise da categoria "Transparência e Explicabilidade" revela convergência global significativa: todas as quatro jurisdições reconhecem estes princípios como condições de legitimidade para sistemas de IA em contextos de alto impacto, como a saúde. Este consenso reflete o debate ético internacional e sinaliza que a "caixa-preta" algorítmica é inaceitável quando decisões afetam direitos fundamentais ou a vida humana.

Contudo, a divergência estratégica reside na natureza jurídica e no propósito atribuídos à transparência. A UE e o Brasil tratam a explicabilidade como direito fundamental do cidadão, consagrado em *hard law*. O GDPR (Artigo 22) e a LGPD (Artigo 20) estabelecem o direito de solicitar revisão de decisões automatizadas e obter informações sobre a lógica subjacente, garantindo ao indivíduo mecanismo legal para contestar decisões. O AI Act europeu aprofunda essa visão ao exigir que sistemas de alto risco sejam "suficientemente transparentes" (Artigo 13) e possuam capacidade de manter registros automáticos (logs) para auditoria (Artigo 12). No Brasil, a RDC 657/2022 exige a descrição de algoritmos nas instruções de uso (Artigo 7º, III), e o PL 2.338/2023 propõe elevar a explicação a direito explícito (Artigo 6º, I).

Os EUA e Singapura, por contraste, concebem a transparência como ferramenta para o bom funcionamento do ecossistema e gestão de riscos, sendo predominantemente guiada por *soft law*. A abordagem não se foca em direito legal à explicação posterior à decisão, mas em mecanismos proativos que capacitam os atores do mercado. O NIST AI RMF defende uma explicabilidade "adequada ao propósito", calibrada com base no risco e no contexto de uso, não imposta uniformemente por lei. A proposta de "rótulos informativos" de Singapura, apresentada no Model AI Governance Framework, fornece informações padronizadas sobre dados de treinamento, desempenho e limitações dos modelos para facilitar a escolha informada. Na saúde, a FDA foca na transparência evolutiva: o PCCP Guidance exige que fabricantes informem continuamente usuários sobre modificações em algoritmos e seus impactos no desempenho. Singapura demonstra distinguir entre transparência (informar ao paciente que interage com IA) e explicabilidade (capacitar o clínico a entender a lógica).

Brasil e Singapura ocupam posições híbridas. O Brasil apresenta trajetória rumo ao modelo europeu garantista, com instrumentos vigentes (RDC 657/2022) impondo obrigações concretas e futuro marco legal consolidando direitos explícitos. Singapura mantém-se no espectro pragmático, propondo instrumentos práticos como rótulos informativos e diferenciação entre públicos-alvo da transparência.

Essa divergência tem implicações práticas. O modelo garantista oferece maior proteção individual e previsibilidade jurídica, mas pode criar barreiras para inovação, especialmente para pequenas empresas que precisam navegar por requisitos legais complexos. O modelo

pragmático favorece agilidade regulatória e inovação, mas depende da maturidade do ecossistema e pode deixar lacunas de proteção onde mecanismos de mercado falham. Enquanto o primeiro protege o cidadão por meio de direitos legalmente exigíveis e fiscalização por agências reguladoras, o segundo busca proteger o ecossistema por meio da promoção de mercado informado e boas práticas voluntárias.

7.2.2 ROBUSTEZ, SEGURANÇA E PROTEÇÃO

A análise desta categoria revela uma convergência unânime: todas as jurisdições identificam a segurança como pré-requisito não negociável para sistemas de IA em saúde. Este consenso reconhece que a confiança pública depende da segurança técnica, proteção contra ataques cibernéticos e capacidade de gerenciar riscos ao longo do tempo. A convergência abrange três dimensões interligadas: segurança do produto, segurança dos dados e cibersegurança.

A divergência reside no "como" essa segurança é alcançada, verificada e garantida. A UE e o Brasil buscam garantir segurança por meio de obrigações legais detalhadas e requisitos técnicos explícitos codificados em *hard law*, partindo da premissa que a segurança deve ser resultado demonstrável de conformidade com padrões predefinidos e verificáveis por autoridades. A UE apresenta abordagem integrada e de diferentes camadas: o MDR estabelece Requisitos Gerais de Segurança e Desempenho rigorosos (Anexo I), o GDPR exige medidas técnicas e organizativas adequadas (Artigo 32) e promove segurança desde a concepção (Artigo 25), e o AI Act prescreve "exatidão, solidez e cibersegurança" (Artigo 15), detalhando medidas contra vulnerabilidades específicas da IA. O Brasil segue filosofia similar: a RDC 657/2022 demanda que software seja desenvolvido considerando princípios do ciclo de vida, gestão de risco e segurança das informações (Artigo 17, III), com conformidade a padrões internacionais (IEC 62304, ISO 14971) comprovada através de Dossiê Técnico. A LGPD complementa ao exigir medidas de segurança desde a concepção (Artigo 46).

Os EUA e Singapura focam na qualidade do processo de desenvolvimento e na maturidade da governança organizacional, partindo da premissa que a segurança não pode ser garantida por lista estática de requisitos, mas de processos robustos de gestão de risco aplicados continuamente. O NIST AI RMF trata "Seguro, Protegido e Resiliente" como características gerenciáveis através de metodologia sistemática que permite às organizações identificar, avaliar e mitigar riscos contextualizadamente. O Executive Order 14110 estabelece que a "IA deve ser segura e protegida" (Seção 2a), delegando ao NIST o desenvolvimento de diretrizes e padrões,

incluindo metodologias de teste cibernético com hackers éticos (Seção 4.1). A abordagem da FDA com o PCCP exemplifica esse modelo: a segurança de algoritmo de aprendizado contínuo é garantida não pela análise regulatória de cada versão, mas pela robustez do "Protocolo de Modificação" que a empresa deve seguir, contendo critérios de aceitação pré-definidos e métodos de validação rigorosos. Singapura promove segurança como capacidade verificável através de mecanismos de mercado, incentivando testagem e auditoria por terceiros e uso de simulação de ataque cibernético. A camada de *hard law* (Health Products Regulations) complementa com sistema reativo de monitoramento e reporte obrigatório de defeitos (Parte VIII).

O modelo de conformidade prescritiva oferece maior clareza regulatória e segurança jurídica para desenvolvedores e usuários, mas pode tornar-se obsoleto face à evolução tecnológica e criar barreiras custosas. O modelo de qualidade processual oferece maior flexibilidade e agilidade para inovação, mas exige alta maturidade de governança organizacional e depende de sistema de supervisão de mercado ou litígio para corrigir falhas. Uma questão emerge: qual modelo é mais adequado para dispositivos médicos com IA de aprendizado contínuo? O modelo prescritivo garante supervisão contínua mas pode retardar a capacidade de melhorar algoritmos em resposta a novos dados clínicos. O modelo processual permite evolução ágil mas transfere maior responsabilidade ao fabricante.

7.2.3 RESPONSABILIDADE (ACCOUNTABILITY)

A análise comparativa dos quatro frameworks regulatórios, sintetizada no Quadro 9, permite identificar padrões de convergência e divergência que transcendem as particularidades nacionais, revelando filosofias regulatórias distintas sobre o papel da IA na sociedade. Esta seção explora essas diferenças estratégicas, categoria por categoria, extraindo arquétipos que podem orientar tanto o desenvolvimento de políticas públicas quanto a compreensão da governança global da IA em saúde. A estrutura de cada subseção segue o mesmo padrão analítico: identificar a convergência, aprofundar a divergência estratégica e apontar os arquétipos que emergem da comparação.

A análise revela convergência conceitual: todas as jurisdições reconhecem que a complexidade técnica e natureza distribuída dos sistemas de IA exigem mecanismos claros para atribuição de deveres e responsabilização por resultados. O princípio de que desenvolvedores, criadores de aplicações, distribuidores e implementadores devem prestar contas por suas ações

e impactos é “universalmente aceito”, refletindo reconhecimento de que a "diluição de responsabilidade" é um dos maiores riscos da IA em saúde.

A UE e o Brasil focam em distribuir e definir obrigações legais explícitas para cada ator ao longo da cadeia de valor, estabelecendo regime de responsabilização jurídica multinível, onde a accountability é compartilhada mas não diluída: deve ser legalmente atribuída e exigível. O AI Act define obrigações dos "prestadores" (Artigo 16), incluindo implementação de sistema de gestão de risco, garantia de qualidade dos dados, elaboração de documentação técnica e estabelecimento de sistemas de registro automático (logs) e especifica obrigações dos "responsáveis pela implantação" (Artigo 26), como realização de avaliações de impacto e garantia de supervisão humana. O GDPR complementa ao estabelecer o "princípio da responsabilidade" (Artigo 5(2)), invertendo o ônus: o controlador não deve apenas cumprir a lei, mas ser capaz de demonstrar cumprimento, operacionalizado pelas Avaliações de Impacto (DPIA). O MDR cria a figura da "Pessoa responsável pela observância da regulamentação" (Artigo 15), estabelecendo ponto focal de responsabilidade profissional dentro de cada fabricante.

O Brasil adota lógica similar. A LGPD consagra o princípio da "responsabilização e prestação de contas" (Artigo 6º, X), exigindo demonstração de adoção de medidas eficazes, com o Relatório de Impacto servindo como instrumento probatório, e estabelece regime de responsabilidade solidária entre controlador e operador (Artigo 42). A RDC 751/2022 atribui responsabilidade primária ao "detentor" do registro e ao "responsável técnico" (Artigo 4º), criando dupla camada de accountability. O PL 2.338/2023 fortalece esse modelo ao prever Capítulo V dedicado à Responsabilidade Civil, com mecanismo de inversão do ônus da prova em favor da vítima (Artigo 37): em caso de dano causado por sistema de alto risco, cabe aos desenvolvedores e aplicadores provarem proativamente a segurança de seus sistemas e ausência de nexo causal, não à vítima provar a falha. Neste modelo, accountability é concebida como questão de direito e responsabilidade civil por danos, com o Estado atuando como árbitro que define em lei os deveres de cada elo da cadeia e garante mecanismos de reparação para as vítimas.

Os EUA e Singapura alocam responsabilidade primariamente dentro da organização que desenvolve ou implementa IA, focando na criação de estruturas de governança interna e cultura organizacional de gestão de risco. O NIST AI RMF dedica sua função "Governar" a este princípio, enfatizando definição de "funções e responsabilidades" claras e estabelecimento de políticas organizacionais para supervisão de gestão de riscos. Responsabilidade é vista como consequência de boa governança interna. O Executive Order 14110 aplica essa lógica ao

governo, determinando que cada agência designe Chief AI Officer (Seção 10), criando ponto focal de accountability interna institucionalizada. Na FDA, a responsabilidade do fabricante é concentrada na fase pré-mercado: o fabricante tem ônus de provar segurança e eficácia de seu produto para obter autorização, mas após aprovação torna-se responsável por seguir o plano submetido. Singapura leva esse modelo prodondo accountability baseada no "nível de controle" de cada ator, formalizada por meio de contratos (SLAs) entre desenvolvedores e implementadores. Neste modelo, accountability é concebida como questão de governança organizacional e autorregulação orientada.

Esta divergência revela duas visões sobre o papel do Estado e do Direito. No modelo europeu-brasileiro, o Estado atua como árbitro que define em lei os deveres de cada elo, estabelecendo mecanismos de supervisão antes da ocorrência de problemas e garantindo reparação após eventos danosos através de regimes de responsabilidade civil robustos. No modelo americano-singapurense, o Estado atua como facilitador que define metodologias de governança que organizações devem internalizar, confiando mais em mecanismos de mercado e contratuais para alocar riscos e responsabilidades. O modelo de responsabilidade em cadeia de valor oferece maior segurança jurídica e proteção ao lesado, mas pode gerar complexidade operacional e custos elevados. O modelo de responsabilidade organizacional incentiva maturidade de governança corporativa e oferece maior flexibilidade, mas pode dificultar a atribuição de culpa em cadeia complexa, deixando vítimas em posição de maior vulnerabilidade jurídica.

7.2.4 JUSTIÇA E PRIVACIDADE

A análise revela que todas as jurisdições utilizam legislação de proteção de dados pessoais como alicerce para governança de IA. Este consenso reflete o reconhecimento de que sistemas de IA são intrinsecamente intensivos em dados e que "dados de saúde" são particularmente sensíveis. Todas possuem instrumentos de *hard law* para proteção de dados: UE com GDPR, EUA com regulações setoriais, Singapura com PDPA e Brasil com LGPD, compartilhando princípios como limitação de finalidade, minimização de dados, transparência e segurança.

Na segunda dimensão, justiça algorítmica (mitigação de vieses e combate à discriminação automatizada), a divergência torna-se evidente. A justiça algorítmica, ao contrário da privacidade, a questão de como prevenir e remediar discriminações causadas ou

amplificadas por algoritmos é relativamente recente e ainda está em formação. A comparação revela três abordagens distintas.

A UE transforma justiça algorítmica em obrigação de engenharia detalhada e prescrita em *hard law*, elevando-a de princípio ético aspiracional a requisito técnico auditável. A privacidade é garantida pelo GDPR, que classifica "dados relativos à saúde" como categoria especial de dados sujeita a proteção reforçada (Artigo 9). Já o AI Act, seu Artigo 10 vai além de princípios genéricos e impõe deveres técnicos: exige que conjuntos de dados de treinamento, validação e teste sejam "pertinentes, suficientemente representativos e tanto quanto possível, isentos de erros e completos", exige exame ativo para detecção de vieses e adoção de medidas para mitigação, transformando o processo de reduzir o impacto de vieses cognitivos e de outros vieses de aspiração ética em etapa obrigatória e documentada. Neste modelo, justiça é etapa obrigatória, auditável e juridicamente exigível, cuja não conformidade pode resultar em sanções.

Os EUA abordam justiça algorítmica como prioridade de política de Estado, enquadrada como questão de aplicação de leis de direitos civis existentes, cuja implementação técnica é delegada às organizações como processo de gestão de riscos voluntário guiado por *soft law*. O Executive Order 14110 declara que a administração "não irá tolerar o uso de IA para desfavorecer" grupos vulneráveis (Seção 2d), direcionando agências a aplicar rigorosamente leis de direitos civis e criando força-tarefa específica do HHS para promover "uso equitativo de IA na saúde" (Seção 8.b). O NIST AI RMF e PCCP Guidance traduzem essa diretriz em *soft law* técnica, tratando "gestão de viés" como risco técnico a ser mapeado, medido e mitigado. A diferença estrutural em relação à UE: ausência de *hard law* específica que detalhe tecnicamente obrigações de redução de vieses e lacuna da lei federal de privacidade abrangente, fragmentando territorialmente a proteção de dados. Justiça algorítmica é prioridade política clara, mas efetivação depende mais de governança organizacional proativa e aplicação posterior de leis antidiscriminação.

Singapura e Brasil enquadram-se em modelo híbrido. Singapura possui lei de privacidade (PDPA), mas trata justiça algorítmica predominantemente como boa prática de engenharia via *soft law* técnico, recomendando "remoção de viés" como questão de qualidade de dados, enfatizando controles de qualidade rigorosos e testes de desempenho em diferentes subgrupos demográficos. O Brasil possui alicerce de privacidade mais forte: a LGPD garante privacidade comparável ao GDPR e estabelece explicitamente princípio da "não discriminação" (Artigo 6º, IX). O PL 2.338/2023 representa salto qualitativo, criando direito explícito à correção de vieses discriminatórios (Artigo 5º, III) e obrigando agentes de alto risco a adotar

"medidas para mitigar e prevenir vieses" (Artigo 18). Contudo, há diferenças: enquanto o AI Act detalha requisitos para *datasets*, o PL 2.338/2023 delega especificação das medidas técnicas concretas à futura autoridade competente. O Brasil está em trajetória rumo ao modelo europeu, mas efetivação plena depende de regulamentação futura que especifique tecnicamente as obrigações.

A abordagem da UE oferece maior proteção antes de problemas ocorrerem e certeza jurídica: desenvolvedores sabem exatamente o que devem fazer, reguladores possuem critérios claros e vítimas têm base legal explícita para contestação, mas pode criar custos de conformidade elevados. A abordagem dos EUA depende da proatividade de organizações, com fiscalização e responsabilização ocorrendo frequentemente após denúncia ou descoberta de dano discriminatório, oferecendo maior flexibilidade, mas podendo deixar lacunas de proteção. Singapura oferece *soft law* técnico sofisticado sem coerção legal. O Brasil mostra trajetória de convergência com modelo europeu, mas efetivação depende de regulamentação futura.

7.2.5 GOVERNANÇA DO CICLO DE VIDA DO ALGORITMO

A análise revela: todas as jurisdições superaram paradigma regulatório estático tradicional, reconhecendo que a natureza dinâmica da IA, especialmente modelos de aprendizado contínuo que melhoram através da exposição a novos dados, exige abordagem de governança estendida por todo ciclo de vida do sistema. Este consenso representa ruptura em relação ao modelo tradicional de dispositivos médicos, onde aprovação regulatória "no dia do lançamento" era suficiente. No contexto da IA em saúde, há reconhecimento explícito que mecanismos ativos e contínuos de supervisão pós-mercado são essenciais para garantir que sistemas que evoluem mantenham segurança, eficácia e ausência de vieses.

A divergência diz respeito aos mecanismos regulatórios escolhidos para operacionalizar essa supervisão contínua: como regular produtos que mudam constantemente? A comparação expõe dois modelos que abordam este dilema de forma distinta.

Os EUA e Singapura enfrentam o desafio através de mecanismo regulatório flexível, pré-acordado e baseado em processos, partindo da premissa que a agência reguladora não deve aprovar produto "travado", mas sim aprovar e validar processo de mudança controlado que permita evolução dentro de parâmetros predefinidos e metodologicamente rigorosos. A inovação central da FDA é o PCCP, onde o fabricante submete para aprovação não apenas o algoritmo inicial, mas plano detalhado especificando escopo das futuras modificações previstas, "Protocolo de Modificação" para validar cada mudança e mecanismos de transparência para

comunicar mudanças aos usuários. Isso funciona como "contrato regulatório prospectivo": a FDA aprova não o futuro estado do algoritmo, mas qualidade e robustez do processo que o fabricante usará para gerenciar sua evolução. Uma vez aprovado o PCCP, enquanto o fabricante seguir o protocolo validado, pode atualizar seu modelo sem necessidade de nova submissão regulatória completa a cada modificação. Governança desloca-se da aprovação produto-por-produto para aprovação processo-por-processo. O NIST AI RMF complementa ao fornecer metodologia de gestão de riscos contínua e interativa. Singapura adota filosofia similar via AIHGle, recomendando "aprendizado em paralelo": modelos novos devem ser treinados e validados "em segundo plano" enquanto modelo aprovado continua operando, implementados apenas após demonstrarem desempenho superior e segurança equivalente ou melhorada.

A UE e Brasil governam ciclo de vida através de regime de monitoramento contínuo obrigatório combinado com gatilhos legais que forcem reavaliação regulatória formal sempre que o sistema sofre mudança ultrapassando determinados limiares, partindo da premissa que supervisão deve ser contínua mas que autoridade reguladora deve manter-se como controladora de acesso em pontos críticos, garantindo controle direto sobre mudanças substanciais. O AI Act implementa essa lógica através de mecanismos: exige Sistema de Gestão de Riscos abrangendo "todo ciclo de vida do sistema de IA" com revisões regulares (Artigo 9), impõe sistema de acompanhamento pós-comercialização coletando dados sobre desempenho (Artigo 72) e define "modificação substancial" como gatilho regulatório (Artigo 3, item 23): alteração que "afete a conformidade do sistema de IA com os requisitos" ou "resulte numa alteração da finalidade prevista" força nova avaliação de conformidade antes de implementação. O Brasil adota filosofia similar: a RDC 751/2022 estabelece governança através de revalidação periódica e necessidade de aprovação prévia para alterações, a RDC 657/2022 estipula que alterações que "afetem significativamente as funcionalidades clínicas, segurança e eficácia" demandam nova análise da Anvisa (Artigo 16), criando controle de mudanças pós-mercado rigoroso. A LGPD governa o ciclo de vida dos dados, impondo "limitação da finalidade" e "limitação da conservação" (Artigos 15 e 16), impactando e limitando o ciclo de vida do algoritmo. O PL 2.338/2023 reforça ao conceber a Avaliação de Impacto Algorítmico como "processo interativo contínuo" (Artigo 26), não como documento estático produzido uma única vez, devendo ser atualizada periodicamente e sempre que houver mudanças significativas.

A divergência tem implicações para a velocidade de inovação. O modelo americano é projetado para permitir rápida implementação de melhorias em algoritmos, promovendo agilidade regulatória e inovação incremental contínua, mas transfere responsabilidade para governança interna do fabricante e exige maturidade técnica e de processos, dependendo de

sistema de supervisão de mercado pós-aprovação e capaz de verificar se protocolos aprovados estão sendo seguidos. O modelo europeu e brasileiro oferece maior segurança regulatória e controle direto pela autoridade competente, garantindo que nenhuma mudança crítica ocorra sem nova validação formal externa e independente, protegendo mais efetivamente contra risco de que fabricante implemente mudanças inadequadamente validadas. No entanto, esse processo pode ser lento, burocrático e custoso, criando potencial desincentivo para que fabricantes realizem melhorias incrementais frequentes, podendo atrasar a chegada de inovações benéficas aos pacientes e reduzir capacidade de resposta ágil a problemas identificados no mundo real. Uma questão emerge: qual modelo é mais adequado para diferentes classes de risco? É possível um modelo híbrido onde dispositivos de risco mais baixo operem sob regime tipo PCCP enquanto dispositivos de risco mais alto mantenham supervisão reativa mais rigorosa?

7.3 SÍNTESE DA ANÁLISE CATEGORIAL

A análise categorial comparativa das cinco dimensões que estruturam a governança ética e regulatória de IA em saúde estabeleceu compreensão das filosofias regulatórias distintas que moldam a governança da IA nas quatro jurisdições estudadas. A jornada através de transparência e explicabilidade, robustez e segurança, responsabilidade, justiça e privacidade, e governança do ciclo de vida do algoritmo revelou padrões consistentes e as particularidades nacionais.

A análise demonstrou que, embora haja convergências conceituais significativas, todas as jurisdições reconhecem a importância desses princípios éticos para garantir sistemas de IA confiáveis em contextos de saúde, as divergências estratégicas e operacionais são estruturais. Essas divergências refletem diferentes visões sobre aspectos fundamentais da governança tecnológica: o papel do Estado (árbitro que define obrigações versus orientador que capacita organizações), a natureza da regulação (prescritiva e detalhada versus processual e metodológica), os instrumentos jurídicos privilegiados (*hard law* vinculativo versus *soft law* orientador) e o equilíbrio desejado entre proteção de direitos fundamentais e promoção da inovação tecnológica.

Dois arquétipos regulatórios principais se apresentam de forma consistente através das cinco categorias analisadas, como manifestações coerentes de filosofias de governança distintas. O primeiro, liderado pela UE e seguido pelo Brasil, privilegia *hard law* tecnicamente detalhado, estabelece obrigações técnicas explícitas e auditáveis para desenvolvedores e implementadores, distribui responsabilidade legal ao longo da cadeia de valor com mecanismos

de reparação de danos, implementa proteção de direitos fundamentais antes de problemas ocorrerem e mantém supervisão regulatória contínua com controle estatal direto em pontos críticos do ciclo de vida dos sistemas. Este arquétipo concebe a regulação como instrumento de proteção de direitos, onde o Estado atua como árbitro que define claramente as regras e garante sua observância através de autoridades reguladoras com poder de fiscalização e sanção.

O segundo arquétipo, liderado pelos EUA e por Singapura, privilegia *soft law* orientador e frameworks voluntários, enfatiza metodologias de gestão de risco que as organizações devem internalizar, foca na governança organizacional interna e na criação de cultura de risco, baseia a accountability na maturidade de processos e capacidade de demonstração de conformidade, e implementa regulação adaptativa que confia mais em mecanismos de mercado (reputação, certificação voluntária, auditorias de terceiros), contratuais e de responsabilização após ocorrência de danos do que em prescrições técnicas estabelecidas previamente. Este arquétipo concebe a regulação como instrumento de capacitação e orientação, onde o Estado atua como facilitador que define princípios e metodologias, confiando na autorregulação orientada e na disciplina de mercado.

Brasil e Singapura ocupam posições híbridas que combinam elementos de ambos os arquétipos, mas com vetores evolutivos opostos. O Brasil demonstra trajetória rumo ao arquétipo garantista-prescritivo europeu: a LGPD estabeleceu base de proteção de dados e não-discriminação; as RDCs 751/2022 e 657/2022 da Anvisa implementaram requisitos técnicos para dispositivos médicos com software; e o PL 2.338/2023, em fase avançada de tramitação, propõe criar direitos explícitos, obrigações técnicas para sistemas de alto risco e mecanismos protetivos robustos. Singapura, mantém-se deliberadamente no espectro pragmático-processual, refinando continuamente seu modelo de *soft law* técnico. Sua abordagem combina lei de privacidade (PDPA) como camada de proteção com guias técnicos elaborados que introduzem conceitos como rótulos informativos para modelos de IA e formalização contratual da divisão de responsabilidades através de SLAs. Esta escolha reflete estratégia nacional de posicionar-se como centro de conexão de inovação em IA com governança ágil e orientada por padrões técnicos.

As implicações dessas diferentes escolhas regulatórias são múltiplas e envolvem decisões que implicam abrir mão de uma coisa para ganhar outra. O arquétipo garantista-prescritivo oferece maior segurança jurídica para todos os atores (desenvolvedores sabem exatamente o que devem fazer, reguladores possuem critérios claros de auditoria, cidadãos têm direitos exigíveis), maior proteção de direitos fundamentais antes de problemas ocorrerem e mecanismos mais robustos de reparação de danos. No entanto, pode criar custos de

conformidade elevados que funcionam como barreiras de entrada (especialmente para *startups* e pequenas empresas), rigidez que dificulta adaptação rápida a mudanças tecnológicas e potencial desincentivo à inovação incremental contínua devido à necessidade de reavaliações regulatórias frequentes.

O arquétipo pragmático-processual oferece maior flexibilidade e agilidade para inovação tecnológica, custos de conformidade potencialmente menores e capacidade de adaptação rápida às mudanças no estado-da-arte técnico. No entanto, depende da maturidade de governança das organizações, pode deixar lacunas de proteção para cidadãos em situações em que mecanismos de mercado falham, coloca maior ônus probatório sobre vítimas de danos e oferece menor previsibilidade jurídica devido à dependência de interpretações caso a caso de princípios gerais.

Permanece uma questão não resolvida: como equilibrar adequadamente a proteção de direitos fundamentais e a segurança do usuário com o incentivo à inovação tecnológica necessária para melhorar os sistemas de saúde? Regimes excessivamente rígidos podem desencorajar o desenvolvimento e adoção de IA em saúde devido ao risco legal e aos custos de conformidade percebidos, potencialmente privando pacientes de avanços benéficos. Regimes excessivamente flexíveis podem criar situações onde os custos sociais de falhas algorítmicas (discriminação, danos à saúde, violações de privacidade) são suportados por vítimas sem recursos adequados de reparação, minando a confiança pública essencial para a legitimidade social da tecnologia.

A experiência comparada dos próximos anos, com a entrada em vigor e implementação prática do AI Act europeu (2025-2027), a possível aprovação e regulamentação do PL 2.338/2023 brasileiro, a evolução do arcabouço americano sob novas administrações e o refinamento contínuo do modelo de Singapura, será importante para avaliar qual modelo, ou que combinação, melhor equilibram estes objetivos em diferentes contextos nacionais e estágios de maturidade tecnológica.

Para o Brasil, país com aspirações de ser protagonista em saúde digital, de acordo com a Estratégia de Saúde Digital para o Brasil, um sistema de saúde pública de dimensões continentais que poderia beneficiar-se de ferramentas de IA para diagnóstico, triagem, alocação de recursos e personalização de tratamentos, a escolha estratégica sobre qual caminho seguir não é meramente técnica ou jurídica, mas política e social: que tipo de ecossistema de IA em saúde o país deseja construir e para quem?

8. DISCUSSÃO

A análise comparativa apresentada no capítulo anterior revelou a existência de dois arquétipos regulatórios distintos que moldam a governança global de IA na saúde: o garantista-prescritivo e o pragmático-processual. Este capítulo aprofunda a discussão desses achados, explorando suas implicações estratégicas para o contexto brasileiro. Inicialmente, apresenta-se uma síntese dos principais achados da pesquisa, consolidando os padrões identificados na análise comparativa. Em seguida, discute-se o posicionamento regulatório do Brasil frente aos modelos internacionais, analisando os desafios e oportunidades que emergem dessa trajetória. A seção subsequente examina as implicações práticas para a gestão de IA nas instituições de saúde, traduzindo os achados regulatórios em orientações operacionais. Baseando-se nessas análises, são então apresentadas recomendações estratégicas para diferentes atores do ecossistema. Por fim, reconhecem-se as limitações metodológicas do estudo e propõe-se uma agenda de pesquisas futuras.

8.1 SÍNTESE DOS ACHADOS DA PESQUISA

A análise comparativa dos quatro *frameworks* regulatórios revelou, primeiramente, convergência conceitual: todas as jurisdições reconhecem os cinco princípios éticos de IA confiável como fundamentais para a governança de sistemas de IA em contextos de alto impacto como a saúde. Este achado corrobora a observação de Jobin, Ienca e Vayena (2019), que identificaram convergência global em torno de princípios como transparência, justiça e responsabilização em sua análise de 84 documentos éticos sobre IA. Da mesma forma, Fjeld et al. (2020) documentaram que, apesar das diferenças culturais e políticas, existe notável alinhamento internacional sobre os valores fundamentais que devem orientar o desenvolvimento de IA. Não há divergência sobre a importância da transparência, da segurança, da responsabilização, da proteção de dados e da supervisão contínua. Este consenso internacional reflete a maturação do debate ético sobre IA e sinaliza que a questão não é mais "se" estes princípios devem ser observados, mas "como" implementá-los na prática.

Contudo, as divergências estratégicas e operacionais são profundas e estruturais, refletindo visões fundamentalmente distintas sobre o papel do Estado, a natureza da regulação e o equilíbrio entre proteção de direitos e promoção da inovação. Cath et al. (2018) já haviam alertado que, embora haja consenso sobre princípios éticos abstratos, persistem profundos desacordos sobre como operacionalizá-los em sistemas regulatórios concretos. Veale e Brass

(2019) argumentam que essas divergências não são meramente técnicas, mas refletem diferentes tradições jurídicas e filosofias políticas sobre governança tecnológica. Conforme detalhado no Capítulo 7, a análise identificou dois arquétipos regulatórios principais que se manifestaram de forma consistente através das cinco categorias analisadas: o Arquétipo Garantista-Prescritivo e o Arquétipo Pragmático-Processual.

O primeiro, liderado pela UE e seguido em trajetória pelo Brasil, privilegia instrumentos de *hard law* tecnicamente detalhados, estabelece obrigações técnicas explícitas e auditáveis para desenvolvedores e implementadores, distribui responsabilidade legal ao longo da cadeia de valor com mecanismos robustos de reparação de danos, consagra direitos fundamentais em lei (como o direito à explicação e à correção de vieses) e mantém supervisão regulatória contínua com controle estatal direto em pontos críticos do ciclo de vida. Neste modelo, o Estado atua como árbitro que define as regras e garante seu cumprimento através de autoridades reguladoras com poder de fiscalização e sanção.

O segundo, liderado pelos EUA e por Singapura, privilegia *soft law* orientador e *frameworks* voluntários, enfatiza metodologias de gestão de risco que as organizações devem internalizar, foca na governança organizacional interna e na criação de cultura de risco, baseia a *accountability* na maturidade de processos e capacidade de demonstração de conformidade, e implementa regulação adaptativa que confia em mecanismos de mercado, contratuais e de responsabilização após ocorrência de danos. Neste modelo, o Estado atua como facilitador que define princípios e metodologias, confiando na autorregulação orientada e na disciplina de mercado.

O Brasil demonstra trajetória rumo ao arquétipo garantista-prescritivo europeu: a LGPD estabeleceu base sólida de proteção de dados; as RDCs 751/2022 e 657/2022 da Anvisa implementaram requisitos técnicos rigorosos; e o PL 2.338/2023, em fase final de tramitação, propõe criar arcabouço completo e protetivo, incluindo direitos explícitos, obrigações técnicas detalhadas para sistemas de alto risco e mecanismos como inversão do ônus da prova em casos de dano. Singapura, mantém-se deliberadamente no espectro pragmático-processual, refinando continuamente seu modelo de *soft law* técnico como estratégia de posicionamento como hub de inovação em IA.

Cada arquétipo apresenta vantagens e desvantagens que envolvem trocas inevitáveis. O modelo garantista-prescritivo oferece maior segurança jurídica, proteção robusta de direitos fundamentais e mecanismos claros de reparação de danos, mas pode criar custos de conformidade elevados, barreiras de entrada para inovação e rigidez que dificultam adaptação rápida a mudanças tecnológicas. O modelo pragmático-processual oferece maior flexibilidade,

agilidade para inovação e custos de conformidade menores, mas depende da maturidade de governança das organizações, pode deixar lacunas de proteção onde mecanismos de mercado falham e oferece menor previsibilidade jurídica.

8.2 O POSICIONAMENTO REGULATÓRIO DO BRASIL

O Brasil encontra-se em um momento estratégico de definição de seu modelo regulatório para IA. A análise comparativa demonstra que o país já possui um alicerce regulatório avançado em comparação com muitas nações:

A LGPD estabeleceu uma base de proteção de dados que é estruturalmente similar ao GDPR europeu, classificando dados de saúde como sensíveis e consagrando princípios como transparência, segurança, responsabilização e não-discriminação. Esta lei coloca o Brasil entre as poucas jurisdições com legislação abrangente de proteção de dados, uma lacuna crítica nos EUA.

A regulação setorial da Anvisa, através das RDCs 751/2022 e 657/2022, demonstra maturidade técnica ao estabelecer requisitos específicos para dispositivos médicos com software e IA. A RDC 657/2022 por exige conformidade com padrões internacionais de engenharia de software e gestão de risco, a descrição de algoritmos nas instruções de uso e controle rigoroso de mudanças pós-mercado através do conceito de "alterações significativas". Esta regulação setorial alinha o Brasil às melhores práticas internacionais.

O PL 2.338/2023, em sua versão em tramitação no Senado Federal, representa o consenso político sobre regulação horizontal de IA no país. O projeto incorpora conceitos do debate global: define sistemas de IA de alto risco, estabelece direitos explícitos (como o direito à explicação e à correção de vieses), impõe obrigações técnicas aos desenvolvedores e aplicadores, cria a Avaliação de Impacto Algorítmico como instrumento de governança contínua e propõe a inversão do ônus da prova em favor da vítima em casos de dano.

A trajetória brasileira é convergente com o modelo europeu garantista-prescritivo. Esta escolha estratégica não é acidental, mas reflete valores constitucionais e prioridades de política pública: o Brasil, como país com profundas desigualdades sociais e um sistema de saúde pública universal que atende centenas de milhões de pessoas, opta por um modelo que prioriza a proteção de direitos fundamentais, a redução de vieses discriminatórios e mecanismos robustos de responsabilização, mesmo que isso implique maior complexidade regulatória e potenciais custos de conformidade mais elevados.

Contudo, o Brasil enfrenta desafios críticos que distinguem seu contexto do europeu. Primeiro, a fragmentação institucional: diferentemente da UE, onde regulamentos têm aplicação direta e uniforme, o Brasil possui múltiplas autoridades (Anvisa para dispositivos médicos, ANPD para dados pessoais, futura autoridade de IA) cuja coordenação precisa ser consolidada. Segundo, a capacidade de cumprimento: implementar um modelo prescritivo exige autoridades reguladoras com recursos humanos, técnicos e financeiros robustos para fiscalização e auditoria, um desafio em contexto de restrição orçamentária. Terceiro, a assimetria de maturidade: enquanto na Europa há um ecossistema maduro de empresas com capacidade de conformidade, no Brasil coexistem grandes corporações com recursos e startups e pequenas empresas que podem ter dificuldade em navegar requisitos complexos, criando risco de que a regulação funcione como barreira de entrada inadvertida.

8.3 IMPLICAÇÕES PARA A GESTÃO EM SAÚDE

Para gestores de instituições de saúde no Brasil, a análise comparativa dos *frameworks* internacionais e o posicionamento do país no espectro garantista-prescritivo têm implicações práticas diretas que demandam atenção estratégica.

A conformidade regulatória torna-se competência estratégica essencial. O cenário de "soberania dupla" (Anvisa e LGPD) será ampliado com o futuro marco legal de IA (PL 2.338/2023), exigindo navegação por arcabouço que abrange dispositivos médicos, proteção de dados e governança algorítmica. A não conformidade pode resultar em sanções administrativas, responsabilização civil e danos reputacionais graves.

A decisão de aquisição de IA exige critérios técnicos e jurídicos robustos além de acurácia clínica e custo-benefício. É necessário verificar: documentação técnica comprovando conformidade com padrões de engenharia de *software*; evidência de gestão de viés nos dados de treinamento; explicabilidade adequada ao contexto; mecanismos de rastreabilidade (*logs*); plano de governança do ciclo de vida; e clareza contratual sobre divisão de responsabilidades. A experiência de Singapura com SLAs formalizando responsabilidades entre desenvolvedor e implementador oferece lição valiosa adaptável ao contexto brasileiro.

A governança interna de IA torna-se imperativa, independentemente do porte institucional. Morley et al. (2020) argumentam que, mesmo em contextos de regulação flexível, a governança organizacional robusta é pilar central de conformidade e gestão de riscos em IAa. No Brasil, com regulação prescritiva, isso é condição de sobrevivência. Instituições precisarão: designar responsáveis pela governança de IA; implementar processos formais de avaliação de

risco; estabelecer supervisão humana sobre decisões algorítmicas clínicas; criar protocolos de monitoramento pós-implementação; e manter documentação auditável.

A complexidade dos *frameworks* e a natureza técnica da IA criam lacuna de competências que deve ser endereçada através de conhecimento sobre conceitos fundamentais de IA, princípios de proteção de dados, metodologias de gestão de risco algorítmico, padrões técnicos internacionais e interpretação de obrigações legais. Esta literacia institucional pode ser desenvolvida via capacitação interna, contratação de expertise ou parcerias acadêmicas.

Por fim, a postura em relação à inovação precisa equilibrar oportunidade e prudência. O modelo brasileiro, ao convergir com o europeu garantista-prescritivo, sinaliza prioridade de proteção robusta, mesmo implicando processos mais longos. Gestores devem calibrar expectativas: a adoção de IA em saúde será provavelmente mais lenta que em contextos pró-inovação, mas oferecerá maior segurança jurídica. A estratégia adequada é a adoção progressiva e criteriosa, priorizando casos de uso com benefício clínico claro e capacidade institucional de garantir conformidade e supervisão.

8.4 RECOMENDAÇÕES

Com base nos achados da análise comparativa e nas melhores práticas identificadas nos *frameworks* internacionais, apresentam-se recomendações estratégicas para três grupos de atores do ecossistema de IA em saúde no Brasil.

Para gestores de instituições de saúde, é fundamental estabelecer estrutura formal de governança de IA com designação clara de responsáveis, processos documentados de avaliação de risco e supervisão contínua, inspirando-se no conceito de Chief AI Officer e na função "Governar" do NIST AI RMF. A decisão de aquisição de tecnologia deve incorporar critérios robustos de investigação aprofundada técnica e jurídica, incluindo *checklist* de conformidade com LGPD, RDCs da Anvisa e futuro marco legal, exigindo de fornecedores documentação técnica detalhada, evidências de gestão de viés e planos de governança do ciclo de vida. O investimento em capacitação de equipes sobre literacia regulatória e técnica em IA é prioritário. A implementação de mecanismos de monitoramento pós-implementação com revisões periódicas de desempenho, detecção de degradação e identificação de vieses emergentes deve ser acompanhada de protocolos claros para a suspensão de sistemas problemáticos. A estratégia de adoção deve ser progressiva e criteriosa, priorizando casos de uso com benefício clínico comprovado, risco gerenciável e capacidade institucional de supervisão, evitando a tentação de adoção apressada de tecnologias imaturas.

Para formuladores de políticas públicas, a prioridade deve ser assegurar harmonização e coordenação institucional entre as autoridades envolvidas (Anvisa, ANPD, futura autoridade de IA), criando mecanismos formais de articulação, sendo a experiência europeia com o Guia AIB/MDCG um modelo a considerar. A avaliação de mecanismos de regulação adaptativa para equilibrar proteção e inovação, inspirados no PCCP da FDA, deve contemplar processos de aprovação prospectiva para sistemas de menor risco onde se valida o processo de mudança ao invés de cada versão do produto. O fortalecimento institucional das autoridades reguladoras exige investimento em capacitação, contratação de especialistas, ferramentas de auditoria e treinamento contínuo. A promoção de infraestrutura pública de apoio à conformidade deve incluir bancos de dados representativos da população brasileira e guias técnicos detalhados que auxiliem desenvolvedores e instituições.

O fomento à pesquisa e desenvolvimento em IA para saúde pública, com ênfase no SUS, deve priorizar soluções que incorporem desde o *design* requisitos de conformidade, representatividade populacional e mitigação de vieses. O estabelecimento de mecanismos de monitoramento da efetividade do marco regulatório com indicadores claros permitirá avaliar se os objetivos estão sendo atingidos sem criar barreiras desproporcionais à inovação.

Para reguladores, a ação prioritária deve ser o desenvolvimento de guias técnicos detalhados que traduzam obrigações legais em requisitos operacionais claros, inspirando-se no nível de especificidade do Guia AIB/MDCG europeu e dos guias da FDA. A promoção de diálogo contínuo com o ecossistema através de fóruns permanentes de consulta é essencial para que a regulação evolua informada pelas experiências práticas de implementação e pelos avanços do estado-da-arte técnico. Por fim, o estabelecimento de mecanismos de vigilância pós-mercado robustos e ágeis, com canais eficientes para reporte de incidentes, eventos adversos e suspeitas de viés ou discriminação algorítmica, deve ser complementado por capacidade de investigação rápida e, quando necessário, de suspensão temporária de sistemas que apresentem riscos à segurança do paciente ou violação de direitos.

8.5 LIMITAÇÕES DO ESTUDO

A análise baseou-se exclusivamente em fontes primárias documentais (leis, regulamentos, guias oficiais), sem incluir dados empíricos sobre implementação prática. Não foram realizadas entrevistas com reguladores, gestores ou desenvolvedores, limitando a pesquisa à dimensão normativa sem capturar a dimensão empírica. O recorte intencional de quatro jurisdições (UE, EUA, Singapura e Brasil), embora justificado por sua

representatividade de diferentes arquétipos, exclui outras experiências relevantes como Reino Unido, Canadá, China, Japão e Coreia do Sul.

O recorte temporal (2024-2025) captura um momento específico de um campo em rápida evolução: o AI Act europeu entrará em vigor progressivamente entre 2025-2027, o PL 2.338/2023 brasileiro ainda está em tramitação, e os EUA podem desenvolver novas regulações, tornando as análises potencialmente desatualizadas. O foco específico em saúde, embora necessário para aprofundamento, limita a generalização das conclusões para outros setores onde os frameworks também se aplicam.

Parte significativa da análise do caso brasileiro baseia-se no PL 2.338/2023, ainda em tramitação, conferindo caráter prospectivo às conclusões sobre o "futuro framework brasileiro" que pode não se concretizar exatamente como descrito. Por fim, este trabalho não realizou avaliação de custo-benefício, não estimou custos de conformidade nem analisou impactos econômicos sobre inovação, baseando suas conclusões em análise qualitativa dos instrumentos normativos, não em evidências econômicas empíricas.

8.6 PESQUISAS FUTURAS

As limitações identificadas e os achados desta pesquisa abrem múltiplas oportunidades para investigações futuras que podem aprofundar e complementar a compreensão sobre governança de IA em saúde no Brasil e no cenário internacional.

Estudos empíricos sobre implementação prática são prioritários, realizando pesquisas qualitativas e quantitativas com reguladores, gestores, desenvolvedores e profissionais clínicos para compreender como os frameworks funcionam na prática: principais desafios de conformidade, custos reais, estruturação de governança interna, mecanismos efetivos de supervisão e diferenças entre instituições públicas, privadas e de diferentes portes.

Avaliação de efetividade dos modelos regulatórios após entrada em vigor do AI Act europeu (2025-2027) e eventual aprovação do PL 2.338/2023, através de estudos longitudinais avaliando: número de sistemas aprovados, tempo médio de aprovação, incidentes reportados, percepção de confiança pública e impactos sobre inovação.

Estudos de caso de instituições brasileiras implementando IA podem revelar como navegam o arcabouço regulatório, critérios de seleção de fornecedores, estruturação de governança, competências necessárias e resultados clínicos, operacionais e financeiros. Análise econômica comparativa dos modelos pode estimar custos de conformidade, avaliar impactos sobre investimento em P&D e comparar barreiras de entrada.

Investigação sobre competências necessárias para gestores pode identificar e validar conhecimentos, habilidades e atitudes essenciais, subsidiando desenvolvimento de currículos e programas de capacitação. Expansão da pesquisa comparativa para incluir Reino Unido, Canadá, China, Japão e países latino-americanos ampliaria o espectro de lições disponíveis.

9. CONCLUSÃO

Este trabalho propôs-se a analisar comparativamente os frameworks regulatórios de inteligência artificial na saúde dos EUA, UE e Singapura, buscando identificar tendências, desafios e lições estratégicas aplicáveis ao cenário brasileiro. Através de uma pesquisa documental comparativa estruturada em cinco categorias analíticas derivadas dos princípios de IA confiável da OECD, foi possível mapear não apenas as convergências conceituais que unem o debate global, mas sobretudo as profundas divergências estratégicas que revelam visões distintas sobre o papel do Estado, a natureza da regulação e o equilíbrio entre proteção de direitos fundamentais e promoção da inovação tecnológica. A jornada analítica demonstrou que, embora haja consenso internacional sobre a importância da transparência, segurança, responsabilização, privacidade e supervisão contínua de sistemas de IA, as formas de operacionalizar esses princípios variam substancialmente, cristalizando-se em dois arquétipos regulatórios principais que coexistem e competem no cenário internacional.

A principal contribuição desta pesquisa reside na identificação e caracterização de dois arquétipos regulatórios distintos que estruturam o debate global sobre governança de IA na saúde, oferecendo uma lente analítica que permite compreender as escolhas estratégicas de diferentes jurisdições não como modelos isolados, mas como manifestações de filosofias políticas coerentes sobre o papel do Estado, a natureza da regulação e o equilíbrio entre proteção e inovação. Ao posicionar o Brasil claramente no espectro garantista-prescritivo e demonstrar que essa trajetória reflete escolhas políticas deliberadas alinhadas aos valores constitucionais de um país com profundas desigualdades sociais e um sistema de saúde pública universal, o trabalho transcende a descrição jurídica para oferecer uma análise estratégica das implicações dessa escolha, evidenciando tanto as oportunidades quanto os desafios que emergem de um modelo que prioriza proteção robusta de direitos fundamentais em contexto de capacidades institucionais assimétricas e recursos limitados. A análise demonstrou que não existe um modelo único superior, mas sim escolhas que refletem dilemas entre segurança jurídica e flexibilidade regulatória, entre prevenção e responsabilização, entre prescrição detalhada e

orientação principiológica, cabendo ao Brasil construir um caminho próprio que equilibre esses objetivos com suas especificidades contextuais.

A relevância desta pesquisa situa-se precisamente no momento estratégico que o Brasil atravessa. Com a iminente aprovação de um marco legal específico para IA e a crescente adoção dessas tecnologias no setor da saúde, gestores de instituições, formuladores de políticas públicas e reguladores enfrentam decisões críticas sem dispor de referências claras sobre as melhores práticas internacionais. Este trabalho buscou preencher essa lacuna ao traduzir experiências regulatórias globais em conhecimento estratégico aplicável à realidade brasileira, oferecendo um quadro analítico que permite antecipar desafios, compreender dilemas inevitáveis e tomar decisões mais informadas sobre aquisição, implementação e governança de tecnologias de IA. A análise comparativa demonstrou que não existe um modelo único superior, mas sim escolhas que refletem valores, prioridades e capacidades institucionais distintas, cabendo ao Brasil construir um caminho próprio que equilibre a proteção robusta de direitos com viabilidade econômica e capacidade de implementação.

Este trabalho alinha-se diretamente aos Objetivos de Desenvolvimento Sustentável da Agenda 2030 das Nações Unidas (UNITED NATIONS, 2015), particularmente aos ODS 3 (Saúde e Bem-Estar), 10 (Redução das Desigualdades) e 16 (Paz, Justiça e Instituições Eficazes). Ao analisar frameworks regulatórios que buscam garantir que sistemas de inteligência artificial na saúde sejam transparentes, seguros e livres de vieses discriminatórios, a pesquisa contribui para a Meta 3.8 de alcançar cobertura universal de saúde através de tecnologias acessíveis e confiáveis, para a Meta 10.3 de assegurar igualdade de oportunidades mediante a eliminação de práticas discriminatórias algorítmicas, e para a Meta 16.6 de desenvolver instituições eficazes e responsáveis através de mecanismos robustos de governança de IA. A escolha do Brasil por um modelo garantista-prescritivo que prioriza proteção de direitos fundamentais e mitigação de vieses reflete compromisso com desenvolvimento tecnológico inclusivo e equitativo, essencial para que os benefícios da inteligência artificial na saúde sejam amplamente distribuídos e não aprofundem desigualdades estruturais existentes.

A experiência internacional oferece lições valiosas, mas não fornece respostas prontas. O Brasil, ao optar por um modelo garantista-prescritivo, assume o compromisso de construir um ecossistema de IA na saúde que seja não apenas tecnologicamente avançado, mas fundamentalmente justo, transparente e responsável. Este compromisso exige investimento contínuo no fortalecimento institucional das autoridades reguladoras, na capacitação de gestores e profissionais, na promoção de literacia regulatória e técnica, e no desenvolvimento de infraestrutura pública que facilite a conformidade sem sufocar a inovação. A construção

desse ecossistema demandará esforço coletivo de múltiplos atores, paciência para ajustes iterativos à medida que a tecnologia e a sociedade evoluem, e sobretudo clareza sobre os valores fundamentais que devem orientar o desenvolvimento e a aplicação da IA a serviço da saúde de todos os brasileiros.

10. REFERÊNCIAS

ARRIETA, A. B. et al. **Explainable Artificial Intelligence (XAI):** Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, v. 58, p. 82-115, 2020.

BARDIN, L. (2011). **Análise de conteúdo**. São Paulo: Edições 70.

BATOOL, A.; ZOWGHI, D.; BANO, M. AI governance: a systematic literature review. **AI and Ethics**, v. 5, 14 jan. 2025.

BIONI, Bruno Ricardo. **Proteção de dados pessoais: a função e os limites do consentimento**. Rio de Janeiro: Forense, 2019.

BOŽIĆ, V. **Understanding the Lifecycle of AI Applications: A Comprehensive Analysis**. Preprint, Aug. 2024. Disponível em: https://www.researchgate.net/publication/382888386_Understanding_the_Lifecycle_of_AI_Applications_A_Comprehensive_Analysis. Acesso em: 2 out. 2025.

BRASIL. Congresso Nacional. Senado Federal. **Projeto de Lei nº 2338**, de 2023. Dispõe sobre o uso da Inteligência Artificial. Autoria: Senador Rodrigo Pacheco (PSD/MG). Brasília, DF, 2023. Disponível em: <https://www12.senado.leg.br/radio/1/noticia/2018/10/06/entenda-a-tramitacao-de-um-projeto-de-lei-no-congresso-nacional>. Acesso em: 2 out. 2025

BRASIL. **Lei nº 13.709**, de 14 de agosto de 2018. Lei Geral de Proteção de Dados Pessoais (LGPD). Brasília, DF: Presidência da República, 2018. Disponível em: http://www.planalto.gov.br/ccivil_03/_ato2015-2018/2018/lei/113709.htm. Acesso em: 7 out. 2025.

BRASIL. Senado Federal. **Parecer nº 208, de 2024 – PLEN/SF**, sobre o Projeto de Lei nº 2.338, de 2023. Redação para o turno suplementar do Projeto de Lei nº 2.338, de 2023, do

Senador Rodrigo Pacheco, nos termos da Emenda nº 199 – CTIA (Substitutivo). Brasília, DF: Senado Federal, 10 dez. 2024. Disponível em: https://legis.senado.leg.br/sdleg-getter/documento?dm=9865609&ts=1734645094578&rendition_principal=S. Acesso em: 2 out. 2025.

BRADFORD, A. **The Brussels Effect: How the European Union Rules the World**. [s.l.] Oxford University Press, 2019.

CASTELVECCHI, D. **Can we open the black box of AI?** Nature, v. 538, n. 7623, p. 20–23, 5 out. 2016.

CATH, C. Governing artificial intelligence: Ethical, legal and technical opportunities and challenges. **Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences**, v. 376, n. 2133, 15 out. 2018.

CELLARD, A. A Análise Documental. In: POUPART, J. et al. (Orgs.). **A pesquisa qualitativa: enfoques epistemológicos e metodológicos**. Petrópolis, RJ: Vozes, 2008. p. 295-316.

CIRILLO, D. et al. **Sex and gender differences and biases in artificial intelligence for biomedicine and healthcare**. npj Digital Medicine, v. 3, n. 1, 1 jun. 2020.

EISENHARDT, K. M. (1989). **Building theories from case study research**. Academy of management Review, 14(4), 532-550.

ESTEVA, A. et al. **Dermatologist-level Classification of Skin Cancer with Deep Neural Networks**. Nature, v. 542, n. 7639, p. 115–118, 25 jan. 2017.

EUROPEAN UNION. **REGULATION (EU) 2016/679 OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL**. Disponível em: <https://eur-lex.europa.eu/eli/reg/2016/679/oj/eng>. Acesso em: 2 out. 2025

EUROPEAN UNION. **Regulation - EU - 2024/1689 - EN - EUR-Lex**. Disponível em: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>. Acesso em: 2 out. 2025

FJELD, J. et al. **Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-Based Approaches to Principles for AI**. SSRN Electronic Journal, v. 2020, n. 1, 2020.

GERKE, S.; MINSEN, T.; COHEN, G. **Ethical and Legal Challenges of Artificial intelligence-driven Healthcare**. *Artificial Intelligence in Healthcare*, v. 1, n. 1, p. 295–336, 26 jun. 2020.

GIL, Antonio Carlos. **Como Elaborar Projetos de Pesquisa**. 4. ed. São Paulo: Atlas, 2002.

JIANG, F. et al. **Artificial Intelligence in Healthcare: Past, Present and Future**. *Stroke and Vascular Neurology*, v. 2, n. 4, p. 230–243, 21 jun. 2017.

JOBIN, A.; IENCA, M.; VAYENA, E. The Global Landscape of AI Ethics Guidelines. **Nature Machine Intelligence**, v. 1, n. 9, p. 389–399, 2 set. 2019.

KAMINSKI, M. E. **The Right to Explanation, Explained**. Disponível em: <https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3196985>. Acesso em: 2 out. 2025

LÜDKE, M.; ANDRÉ, M. E. D. A. **Pesquisa em educação: abordagens qualitativas**. São Paulo: EPU, 1986.

MORLEY, J.; FLORIDI, L. **The Ethics of AI in Health Care: An Updated Mapping Review**. 1 jan. 2024.

OBERMEYER, Z. et al. **Dissecting Racial Bias in an Algorithm Used to Manage the Health of Populations**. *Science*, v. 366, n. 6464, p. 447–453, 25 out. 2019.

OECD. **Recommendation of the Council on Artificial Intelligence**. Disponível em: <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>. Acesso em: 2 out. 2025.

OECD. **The OECD Artificial Intelligence (AI) Principles**. Disponível em: <https://oecd.ai/en/ai-principles>. Acesso em: 2 out. 2025.

ORGANIZAÇÃO MUNDIAL DA SAÚDE. **Ethics and governance of artificial intelligence for health**: WHO guidance. Genebra: OMS, 2021.

PALANIAPPAN, K.; YAN, E.; VOGEL, S. **Global Regulatory Frameworks for the Use of Artificial Intelligence (AI) in the Healthcare Services Sector**. Healthcare, v. 12, n. 5, p. 562–562, 28 fev. 2024.

PRICE, W. N.; COHEN, I. G. **Privacy in the Age of Medical Big Data**. Nature Medicine, v. 25, n. 1, p. 37–43, 2019.

PRICE, W. N.; GERKE, S.; COHEN, I. G. **Potential Liability for Physicians Using Artificial Intelligence**. JAMA, v. 322, n. 18, p. 1765–1766, 4 out. 2019.

SHAFFER, G.; POLLACK, M. A. **Hard vs. Soft Law: Alternatives, Complements and Antagonists in International Governance**. Disponível em: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=1426123. Acesso em: 2 out. 2025

UNITED NATIONS. **Transforming our world: The 2030 Agenda for Sustainable Development**. Disponível em: <https://sdgs.un.org/2030agenda>. Acesso em: 5 nov. 2025.

VEALE, M.; BRASS, I. **Administration by Algorithm? Public Management meets Public Sector Machine Learning**. Disponível em: https://www.researchgate.net/publication/332555598_Administration_by_Algorithm_Public_Management_meets_Public_Sector_Machine_Learning. Acesso em: 5 nov. 2025.

VOLLMER, S. et al. **Machine learning and artificial intelligence research for patient benefit: 20 critical questions on transparency, replicability, ethics, and effectiveness**. BMJ, v. 368, 20 mar. 2020.

WHO. **Ethics and governance of artificial intelligence for health**. Disponível em: <https://www.who.int/publications/i/item/9789240029200>. Acesso em: 5 nov. 2025.

WIENS, J. et al. **Do no harm: a roadmap for responsible machine learning for health care.** Nature Medicine, v. 25, n. 9, p. 1337–1340, 19 ago. 2019.

ZHANG XINRUI. **Application of artificial intelligence technology in hospital management systems.** Applied and computational engineering, v. 50, n. 1, p. 229–233, 25 mar. 2024.

DECLARAÇÃO SOBRE USO DE INTELIGÊNCIA ARTIFICIAL

Este trabalho contou com o suporte da ferramenta Claude (Anthropic) para revisão textual, verificação de consistência estrutural e aprimoramento da clareza acadêmica. A IA foi utilizada exclusivamente como assistente de redação, sendo todo o conteúdo, análise crítica, interpretações e conclusões de autoria própria do pesquisador. O uso da ferramenta seguiu princípios éticos de transparência acadêmica, conforme recomendações emergentes sobre o uso responsável de IA em pesquisa científica.