

International Consensus Framework for Ethical Integration of Patient-Reported Outcomes and Digital Biomarkers in AI Healthcare Systems

Joana Seringa, João V. Cordeiro, Rui Santana, Teresa Magalhães

Submitted to: Journal of Medical Internet Research
on: August 24, 2025

Disclaimer: © The authors. All rights reserved. This is a privileged document currently under peer-review/community review. Authors have provided JMIR Publications with an exclusive license to publish this preprint on its website for review purposes only. While the final peer-reviewed paper may be licensed under a CC BY license on publication, at this stage authors and publisher expressly prohibit redistribution of this draft paper other than for review purposes.

Table of Contents

Original Manuscript.....	5
Supplementary Files.....	49

Preprint
JMIR Publications

International Consensus Framework for Ethical Integration of Patient-Reported Outcomes and Digital Biomarkers in AI Healthcare Systems

Joana Seringa¹; João V. Cordeiro²; Rui Santana¹; Teresa Magalhães¹

¹ NOVA National School of Public Health, Public Health Research Centre, Comprehensive Health Research Center, CHRC, REAL, CCAL, NOVA University Lisboa, Lisbon, Portugal Lisboa PT

² Interdisciplinary Centre of Social Sciences, NOVA University Lisbon, Lisbon, Portugal Lisboa PT

Corresponding Author:

Joana Seringa

NOVA National School of Public Health, Public Health Research Centre, Comprehensive Health Research Center, CHRC, REAL, CCAL, NOVA University Lisboa, Lisbon, Portugal

Avenida Padre Cruz

Lisboa

PT

Abstract

Background: Alongside expected benefits, several ethical concerns arise from Artificial Intelligence (AI) based models. From the design to the implementation and subsequent evaluation, it is crucial to map potential ethical concerns regarding the use of AI models in healthcare. Patient-Reported Outcomes (PROs) and Digital Biomarkers (DBs) are being increasingly collected to improve patient-centered healthcare systems. However, due to the sensitive nature of this data, its processing into AI models may raise ethical concerns that should be considered. While general AI ethics frameworks exist, no international consensus has specifically addressed the unique ethical challenges of integrating PROs and DBs in AI healthcare models.

Objective: This study aims to fill this critical gap by establishing the first international expert consensus on ethical, legal, and social considerations specifically for PROs and DBs integration in AI-based healthcare models.

Methods: A mixed-method study was performed. Initially, a narrative review was performed to identify the main ethical considerations regarding integrating PROs and DBs in AI-based models in healthcare. Then, a modified two-round online Delphi study was conducted to reach a consensus on the recommendations for integrating PROs and DBs in AI-based models, among selected international experts.

Results: The findings of the two complementary components of this study (narrative review and modified Delphi study) were organized around five core ethical principles: autonomy, beneficence, non-maleficence, justice, and transparency and accountability. The modified Delphi study achieved high consensus (?80%) on 57 specific recommendations across these principles. Key recommendations included implementing dynamic consent models, establishing continuous model validation protocols, conducting regular impact assessments, ensuring diverse stakeholder engagement to mitigate biases, and maintaining human oversight within AI systems.

Conclusions: This study provides the first comprehensive, internationally validated ethical framework specifically designed for PROs and DBs integration in AI healthcare systems, filling a critical gap in the literature that has primarily focused on general AI ethics rather than the unique challenges posed by patient-generated health data.

(JMIR Preprints 24/08/2025:82925)

DOI: <https://doi.org/10.2196/preprints.82925>

Preprint Settings

1) Would you like to publish your submitted manuscript as preprint?

✓ **Please make my preprint PDF available to anyone at any time (recommended).**

Please make my preprint PDF available only to logged-in users; I understand that my title and abstract will remain visible to all users.
Only make the preprint title and abstract visible.

No, I do not wish to publish my submitted manuscript as a preprint.

2) If accepted for publication in a JMIR journal, would you like the PDF to be visible to the public?

✓ **Yes, please make my accepted manuscript PDF available to anyone at any time (Recommended).**

Yes, but please make my accepted manuscript PDF available only to logged-in users; I understand that the title and abstract will remain visible to the public.

Yes, but only make the title and abstract visible (see Important note, above). I understand that if I later pay to participate in <https://www.jmir.org/preprint/82925>, the full text will be available to the public.

No. Please do not make my accepted manuscript PDF available to anyone. I understand that if I later pay to participate in <https://www.jmir.org/preprint/82925>, the full text will be available to the public.



Original Manuscript

International Consensus Framework for Ethical Integration of Patient-Reported Outcomes and Digital Biomarkers in AI Healthcare Systems

Joana Seringa^{1*}, João V. Cordeiro^{1,2}, Rui Santana¹, Teresa Magalhães¹

¹ NOVA National School of Public Health, Public Health Research Centre, Comprehensive Health Research Center, CHRC, REAL, CCAL, NOVA University Lisboa, Lisbon, Portugal

²Interdisciplinary Centre of Social Sciences, NOVA University Lisbon, Lisbon, Portugal.

*Corresponding author – jm.seringa@ensp.unl.pt

NOVA National School of Public Health 1600-560 Lisbon, Portugal

ABSTRACT

Background: Alongside expected benefits, several ethical concerns arise from Artificial Intelligence (AI) based models. From the design to the implementation and subsequent evaluation, it is crucial to map potential ethical concerns regarding the use of AI models in healthcare. Patient-Reported Outcomes (PROs) and Digital Biomarkers (DBs) are being increasingly collected to improve patient-centered healthcare systems. However, due to the sensitive nature of this data, its processing into AI models may raise ethical concerns that should be considered. While general AI ethics frameworks exist, no international consensus has specifically addressed the unique ethical challenges of integrating PROs and DBs in AI healthcare models.

Objective: This study aims to fill this critical gap by establishing the first international expert consensus on ethical, legal, and social considerations specifically for PROs and DBs integration in AI-based healthcare models.

Methods: A mixed-method study was performed. Initially, a narrative review was performed to identify the main ethical considerations regarding integrating PROs and DBs in AI-based models in healthcare. Then, a modified two-round online Delphi study was conducted to reach a consensus on the recommendations for integrating PROs and DBs in AI-based models, among selected international experts.

Results: The findings of the two complementary components of this study (narrative review and modified Delphi study) were organized around five core ethical principles: autonomy, beneficence, non-maleficence, justice, and transparency and accountability. The modified Delphi study achieved high consensus ($\geq 80\%$) on 57 specific recommendations across these principles. Key recommendations included implementing dynamic consent models, establishing continuous model validation protocols, conducting regular impact assessments, ensuring diverse stakeholder engagement to mitigate biases, and maintaining human oversight within AI systems.

Conclusion: This study provides the first comprehensive, internationally validated ethical framework specifically designed for PROs and DBs integration in AI healthcare systems, filling a critical gap in the literature that has primarily focused on general AI ethics rather than the unique challenges posed by patient-generated health data.

Keywords: Artificial Intelligence; Patient-generated data; Digital biomarkers; Ethics; Patient-reported outcomes; Narrative Review; Delphi study; International consensus

INTRODUCTION

Artificial Intelligence (AI) based models promise to improve public health practice and healthcare delivery. AI approaches have been used to predict various health risks and recommend informed decision-making, resulting in more efficient and higher-quality healthcare [1, 2]. Nevertheless, adopting AI-based models can potentially elicit significant ethical, legal, and social challenges for individuals, organizations, and societies [3]. If neglected or unaddressed, such challenges may result in mistrust or rejection of AI implementation in healthcare and consequent significant opportunity costs. Therefore, a proactive approach toward responsible innovation in this field is needed [4, 5].

Patient-reported outcomes (PROs) measures refer to self-reported instruments designed to measure health, quality of life, health experiences, and related constructs that come directly from patients [6]. PROs are increasingly used across health services to improve care [7]. As these outcomes focus on measuring the aspects of health that matter to patients, they contribute to more patient-centered healthcare [8]. Beyond individual patient care and the benefits of promoting shared decision-making to prioritize, treat, and monitor health status and well-being, collecting PROs may contribute to tracking health outcomes and comparing them to best practices and benchmarks for quality improvement [9, 10]. They are also key measures to value-based payment models, population health strategies, health policy, and systems research [11].

PROs can be collected in various environments (e.g., at home, at work, or in the hospital) and methods (e.g., via paper-and-pen, telephone, web-based form, wearable devices, or a mobile app). Therefore, PRO collection infrastructure should be adaptable to meet this diversity [11]. In parallel, outcome measures are usually challenging to collect and require adequate resources [12, 13].

Generally, clinical outcomes are recorded in notes, unstructured free text, unstandardized, and variable formats [12]. Digital health innovations provide the opportunity to meet this challenge [12] by improving data structure and accessibility and, if correctly designed and implemented, reducing the burden on health professionals and organizations [8]. In addition, commercially available electronic devices, such as smartphones, wearables, and Internet of Things (IoT) technologies, allow for capturing and analyzing real-time health-related data, including vital signs, physical activity, stress, or fatigue [12, 14, 15]. These outcomes are known as digital biomarkers, whose passive collection depends less on active patient and provider participation. This allows for greater collection consistency and frequency and offers an additional source of patient-centered data [12]. Remote and continuous clinical data collection through DBs can enhance diagnostic, monitoring, and therapeutic precision [16, 17] while contributing to empowering individuals to take an active role in managing

their health by providing access to personalized health information and facilitating early detection and prevention of disease [15].

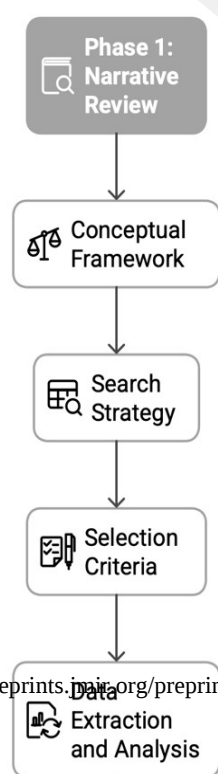
Both PROs and DBs collect patient-generated health-related data, which constitutes sensitive personal information, whose processing (from data collection to data storage) should be responsibly designed to guarantee the protection of individuals' fundamental rights and freedoms [18, 19].

While the ethical implications of AI in healthcare have been comprehensively researched [20–24], the specific challenges posed by integrating patient-generated data, particularly PROs and DBs, in AI models remain underexplored. Unlike general clinical data, PROs and DBs present unique ethical considerations. No previous study has achieved international consensus on ethical frameworks specifically tailored to these data types in AI healthcare applications. Therefore, this paper aims to summarize ethical, legal, and social considerations for integrating PROs and DBs in AI-based models for public health and health care and establish consensus on recommendations for their adequate integration.

These considerations are analyzed through the lens of the ethical principles of autonomy, beneficence, non-maleficence, and justice [25]. In addition to reviewing the topic through this lens, we also discuss transparency and accountability as relevant principles to AI adoption in healthcare [26].

METHODS

Study Design



This research employed a mixed-methods approach conducted in two distinct phases. The first phase consisted of a comprehensive narrative review to systematically identify and analyze ethical considerations related to the integration of PROs and DBs in AI-based models in healthcare. Building on these findings, the second phase implemented a modified Delphi technique to establish expert consensus on recommendations for ethically sound integration of PROs and DBs in AI-based healthcare models. Fig. 1 presents the flowchart of the research methodology.

Phase 1. Narrative review

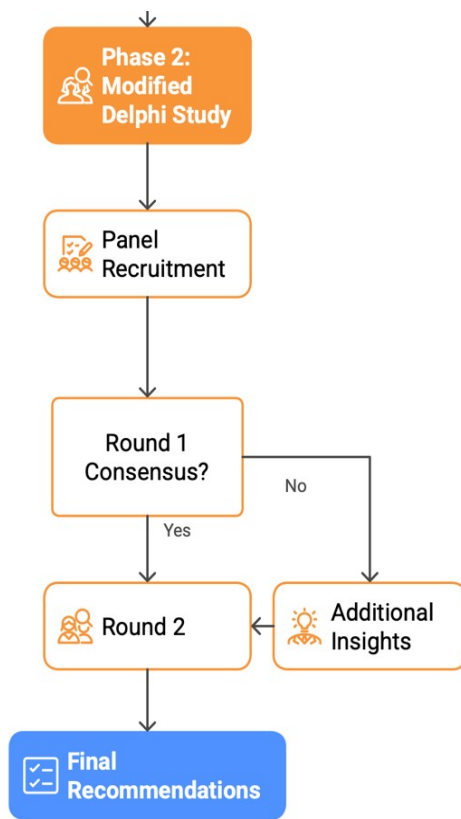


Fig. 1. Research Methodology Flowchart
[Author's elaboration]

No publication date restrictions were applied to capture the evolution of ethical, legal and social discourse on this topic. Articles were initially screened for relevance based on titles and abstracts. Subsequently, a full-text review was conducted for potentially relevant articles. Reference lists of included articles were manually searched through a snowballing technique to identify additional relevant publications.

Data Extraction and Analysis

Data were extracted using a standardized form that captured (1) study characteristics, (2) ethical considerations addressed, and (3) proposed recommendations. Content analysis was performed to identify recurring ethical themes and considerations that informed the development of the Delphi questionnaire.

Phase 2. Modified Delphi Study

Panel Recruitment

An intentional, purposive sampling approach was used to identify 35 experts with diverse expertise in AI, bioethics, clinical outcomes assessment, digital health, and patient advocacy. Informed consent was obtained electronically before participation.

Delphi Process

around the four foundational bioethical principles established by autonomy, non-maleficence, beneficence, and justice. This contemporary ethical considerations of transparency and relevant to AI applications in healthcare [27].

conducted between October 2024 and December 2024 using two is, and Web of Science. To ensure comprehensive coverage, we in various combinations: "Patient-reported Outcomes," "Digital e," "Machine Learning," "Ethics," "Ethical Considerations," l Implications." Boolean operators (AND, OR) were used to icity.

The modified Delphi study was conducted between January 2025 and March 2025 and consisted of two sequential rounds:

Round 1:

- Participants rated their agreement with recommendations derived from the narrative review using a five-point Likert-type scale (1-strongly disagree to 5-strongly agree).
- Open-ended questions allowed experts to provide additional recommendations and contextual insights regarding the processing of PROs and DBs data in AI healthcare models.

Round 2:

- Results from Round 1 were aggregated.
- Items that did not achieve consensus in Round 1 were presented again along with newly identified recommendations from the open-ended responses.
- Participants re-rated their agreement using the same five-point Likert-type scale.

Definition of Consensus

Consensus was predefined as $\geq 80\%$ of participants rating an item as either "agree" (4) or "strongly agree" (5), which aligns with the literature suggesting that an agreement level between 70% to 80% is commonly adopted and regarded as rigorous [28]. Items that did not achieve consensus after the second round were noted but not included in the final recommendations.

Data Analysis

Quantitative data were analyzed using descriptive statistics (percentage agreement) for each recommendation. Qualitative data from open-ended responses were analyzed to identify emergent themes and formulate new recommendations for Round 2.

Ethical considerations and approvals

This study was conducted in accordance with the principles outlined in the Declaration of Helsinki [29]. The research protocol received approval from the NOVA National School of Public Health Research Ethics Committee (reference: CEENSP n° 65/2024).

The Delphi survey was implemented using a fully anonymous data collection methodology. Participants accessed the survey through a public link that did not require login credentials or email addresses. No personal identifying information was collected at any point during the survey process, ensuring complete anonymity of responses.

All participants provided electronic informed consent prior to participation. The anonymous nature of participation was explicitly stated in the consent information.

All aggregated data from the survey rounds were stored in a password-protected Excel document with access restricted to the research team.

No participant received any type of compensation for taking part in this study.

RESULTS

Phase 1: Narrative Review of Ethical, Legal, And Social Considerations of Integrating PROs and DBs in AI-Based Healthcare Models

The narrative review identified key ethical themes organized around five core ethical principles: autonomy, beneficence, non-maleficence, justice, and transparency and accountability. Table 1 presents these principles, the specific themes identified within each principle from the literature review, and the corresponding Delphi items that were developed based on these themes for the consensus-building phase.

Table 1. Ethical Principles, Themes from Literature Review, and Corresponding Delphi Items

Ethical Principle	Themes from Literature Review	Delphi Items
Principle of Autonomy	Data Control and Informed Consent	• Implement granular consent mechanisms, allowing individuals to authorize specific uses of their data.^(*) • Provide individuals control over their data, including the ability to view, manage, modify, and withdraw consent at any time.^(*) • Standardize consent processes across national and European levels to ensure uniformity and ease of patient authorization.

Ethical Principle	Themes from Literature Review	Delphi Items
		• Adopt dynamic and continuously updated consent models, ensuring clear and tailored communication about data collection, usage, and implications. • Periodically review and reaffirm consent mechanisms to ensure sustained autonomy in long-term data collection. • Inform individuals about how their data was used in AI model development, including details on what data was included or excluded. • Use AI to support patient advocacy by empowering individuals to analyze their own data, identify trends, and generate reports for healthcare consultations.
	Effective Communication Tailored to Health Literacy Levels	• Ensure consent processes are comprehensible and adaptable to different levels of health literacy, promoting equitable patient engagement. • Establish transparent protocols for patient communication when AI-driven decisions impact their health, ensuring information is clear, accurate, and precise.
	Transparency and Patient Rights in AI Decision-Making	• Grant patients the "right to explanation," enabling them to understand the rationale behind AI-generated recommendations.^(*) • Share information about the algorithm's performance, the measures taken to identify and prevent errors, and the repercussions for patients' health if the AI algorithm is biased or incorrect.^(*) • Establish mechanisms that allow patients to contest AI-generated insights or decisions they disagree with.^(*)
Principle of Beneficence	Individual and Public Benefits	• Validate AI model effectiveness across different patient groups and stakeholders at multiple checkpoints over time, rather than relying solely on initial assessments.^(*) • Assess AI system benefits from multiple stakeholder perspectives, recognizing distinct priorities among patients, healthcare providers, and administrators.^(*) • Uphold human oversight as a core principle, ensuring AI-driven healthcare decisions are reviewed by qualified professionals.
	Risk-Benefit Assessment	• Implement a comprehensive evaluation framework that tracks performance at each stage of the AI pipeline (data acquisition,

Ethical Principle	Themes from Literature Review	Delphi Items
		preprocessing, modeling, and deployment) to identify sources of deviation and necessary adjustments. • Establish protocols for ongoing assessment and validation of ML models to ensure they continue to deliver beneficial outcomes. • Implement continuous model and data drift monitoring with predefined thresholds to trigger model retraining or resampling when necessary. • Regularly update and recalibrate models to reflect changes in patient populations and healthcare practices. • Support responsible AI development by implementing audits and certifications without stifling innovation. Provide mechanisms for developers to promote responsible AI development, reducing barriers without compromising safety or ethical standards.^(*)
	Interoperability and Data Sharing	• Ensure interoperability between PROs and DBs platforms with other healthcare systems, incorporating real-time data processing capabilities to facilitate timely and informed decision-making for both clinicians and patients.^(*) • Facilitate secure and ethical patient data sharing across organizations to enhance accuracy and reliability in patient care while safeguarding privacy.
Principle of Non-Maleficence	Data Privacy, Security and Compliance	• Safeguard human security through robust cybersecurity measures, preventing unauthorized data access and mitigating potential misuse. • Train staff on cybersecurity best practices, ensuring that individuals involved in the AI lifecycle understand their responsibility to protect patient data from unauthorized access and misuse.^(*) • Apply "security by design" principles, throughout the AI development process, rather than relying solely on policies, to ensure that security measures are inherently built into the technology.^(*) • Integrate privacy and data sovereignty

Ethical Principle	Themes from Literature Review	Delphi Items
		measures into AI solutions, ensuring that patient privacy is maintained without compromising scalability or operational efficiency. • Establish clear notification protocols in case of security breaches or data leaks, ensuring patients are promptly informed. • Establish protocols to address data breaches, including audits to determine the causes and prevent recurrence, and implement more secure procedures to protect patient data. • Implement stringent safeguards to protect stored patient data from unauthorized access. • Perform data and model audits using synthetic data to identify potential vulnerabilities, including security issues or model biases, and ensure that ethical and legal standards are maintained.^(*) • Implement comprehensive audits and oversight systems to ensure ongoing compliance with data protection regulations (e.g., GDPR) and the AI Act, and assess the effectiveness of decisions made based on AI outputs. • Conduct regular impact assessments to evaluate the effects of AI solutions, including patient outcomes (e.g., satisfaction, health improvements), process efficiency, and potential negative consequences (e.g., patient harm, errors). • Ensure AI applications in healthcare undergo Data Protection Impact Assessments (DPIAs) as required by GDPR Article 35.^(*)
	Safety and Error Prevention	• Develop alert mechanisms that flag potential errors or inconsistencies in AI outputs, allowing for timely interventions to prevent harm. • Ensure AI systems provide explainable outputs and maintain human oversight in all healthcare decisions to prevent reliance on AI as a sole decision-maker. Ethical skepticism should guide the integration of AI into decision-making processes.
Principle of Justice	Equity, Inclusion, and Bias Mitigation	• Actively involve diverse populations, including those from varying cultural, socioeconomic, and medical backgrounds, in the design, development, and implementation of digital health tools to ensure equitable access and outcomes.

Ethical Principle	Themes from Literature Review	Delphi Items
		• Ensure that digital health tools are accessible to all populations, particularly vulnerable and underserved communities, with particular attention to mitigating the underuse of these tools by these groups. • Monitor and analyze the demographic profile of populations using digital health tools and take proactive steps to mitigate barriers to access for vulnerable and underserved groups.^(*) • Enable patients to revise their inputs over time, ensuring AI models adapt without introducing bias or undue influence.^(*) • Ensure diverse stakeholder engagement in designing and implementing ML models to mitigate biases. • Ensure training datasets undergo rigorous bias and representativity assessments to achieve fair outcomes across all populations. • Ensure that underrepresented populations are not harmed by biases in ML models, and that measures are in place to regularly detect and correct bias during model development and deployment. • Avoid any form of positive or negative discrimination in feature selection and model design by ensuring that all categories or features used in AI models are balanced, representative, and free from bias.^(*)
	Legal and Ethical Governance	• Implement legal frameworks and harmonize regulations across the EU and member states to create an environment where AI development is both ethical and just, with a focus on ensuring equal access to AI benefits.^(*) • Balance AI regulation to protect patients without stifling innovation, fostering an environment that supports both safety and competitiveness in the EU AI ecosystem.^(*) • Maintain transparency and accountability in AI systems so patients understand how their data is used and trust that it contributes to fair, unbiased outcomes.
Principles of Transparency and Accountability	Transparency and Disclosure	• Implement measures to ensure that AI development and deployment processes, including decision-making criteria and model performance, are well-documented and openly shared with relevant stakeholders.

Ethical Principle	Themes from Literature Review	Delphi Items
		• Document and openly share the AI development and deployment process, including decision-making criteria and model performance. • Regularly publish publicly accessible reports on AI system performance to build trust and demonstrate continuous improvement. • Make information about the databases and statistical methodologies used in AI training and validation publicly available. • Encourage transparent reporting of AI-related errors and unintended consequences by both healthcare providers and AI developers. • Ensure that all stakeholders, including patients and healthcare providers, are informed about any errors or malfunctions that may negatively impact patient care.^(*) • Promptly disclose any issues related to patient data usage, such as breaches or misuse, to maintain trust in AI applications.^(*)
	Interpretability and User Education	• Ensure AI models provide interpretable outputs, including the probability and uncertainty of results, to enhance user understanding. • Enhance the interpretability of ML models to ensure healthcare providers and patients can understand and trust their outputs. • Integrate MLOps best practices in AI-driven healthcare solutions to ensure robust, efficient, and transparent model development, deployment, and maintenance.^(*) • Implement and regularly update training programs for users of AI healthcare tools to ensure they understand the principles of transparency and accountability. • Ensure transparency in AI integration, informing both patients and healthcare professionals of its role and impact in decision-making.
	Human Oversight	• Ensure AI outputs support, rather than dictate, healthcare decisions.^(*) • Ensure human oversight is incorporated into AI systems to guarantee accountability and maintain a role for healthcare providers in decision-making.
	Accountability and Governance	• Establish regular and well-defined audit intervals for AI systems that ensure continuous compliance with ethical and legal standards,

Ethical Principle	Themes from Literature Review	Delphi Items
		while maintaining a balanced approach that upholds accountability without hindering innovation in AI development. • Conduct regular audits and risk assessments to identify and mitigate potential errors and malfunctions. • Establish data lineage protocols to track all data actions, ensuring a non-repudiation method for authorized actors involved in AI healthcare workflows.^(*) • Clearly define the responsibility of regulatory bodies and oversight entities in the continuous evaluation and governance of AI models in healthcare.

Note: The Delphi items were developed based on themes identified in the literature review and were subsequently presented to experts for consensus. Items highlighted with an asterisk ^() were added or modified based on expert suggestions during the first round of the Delphi process.*

Principle of Autonomy

Respect for autonomy is related to respect and support of autonomous decisions [25]. Adopting AI may result in situations where machines are actually or potentially given the responsibility to make important healthcare decisions. According to the autonomy principle, the expansion of AI adoption must respect human autonomy. Furthermore, the notion that AI adoption contributes to patient-centered healthcare suggests that patients should not lose control over their health decisions. In fact, respect for the patient's autonomy consists of involving them in the decision-making process and providing information about the purpose, risks, benefits, and alternatives to an intervention [30].

Therefore, AI systems should assist patients and healthcare professionals in making informed decisions, and human oversight should involve efficient, open monitoring of moral and ethical considerations [22]. This is also the case for AI-based models that process PROs and DBs data to provide valuable insights for informed decision-making.

Respect for autonomy also entails safeguarding confidentiality, privacy, and informed consent [22]. Patients' consent for health data processing by AI algorithms must follow a high standard [31]. Specifically, obtaining adequate informed consent is imperative for collecting DBs, as the continuous collection of passive data can result in a limitation of autonomy due to unanticipated data treatment

[32].

Some authors have already considered that the present informed consent methods for collecting data digitally are unsatisfactory, as patients may not be aware of the type, extent, and potential consequences of the data that is being processed [32]. Consenting to collect data digitally is frequently obtained through near incomprehensible terms-of-service agreements, which may lead to consent without previous and precise acknowledgment, which does not meet adequate ethical standards [33].

It is crucial to involve patients through clear and tailored communication and provide technical training to health professionals [34]. It should also be considered that informed consent comprehension might be influenced by factors such as age, gender, health literacy, and contextual factors [35].

Recommendations to mitigate these concerns should be in accordance with Article 5 of the GDPR - Principles relating to the processing of personal data [36], including the collection of the smallest and simplest data to meet the purposes of data treatment ("data minimization"); the active engagement of individuals through transparent information regarding collected data, the purposes of the data treatment, and the possible intentional and unintentional impact of data treatment on the individual ("lawfulness, fairness, and transparency") [37].

Furthermore, dynamic consent models should be adopted whenever possible, including implementing specific options to select data expiration dates [32]. In addition, ongoing consent models are recommended to ensure continued comprehension and up-to-date voluntariness [38].

Also, patients must be informed that their data will be gathered for health-related purposes while guaranteeing that data can be treated in an aggregated and anonymous manner for public health purposes, including research [34].

Principle of Beneficence

AI adoption in healthcare may result in individual and public benefits. The principle of beneficence is related to relieving or preventing harm while balancing benefits against risks and costs [25].

This principle has been associated with the notion of "well-being" in the Montréal Declaration for a responsible development of AI [39], which states that AI systems must promote the well-being of all individuals.

Beyond the individual and public benefit, the beneficence principle also requires strong evidence that the benefits will exceed the risks [40] in individuals and communities [31].

The risk-benefit ratio assessment evaluates the benefits of AI adoption and its potential harms. This assessment should be based on a human-centered design approach, whereupon the people for whom the algorithm (or other digital solution) is being developed are involved and guide the design process from the beginning and to whom the final solution responds [31].

Beyond the initial assessment of the risk-benefit ratio, the AI algorithm's continuous learning characteristic requires continuous evaluation of whether AI adoption results in patient benefit [41] and whether such benefit outweighs the risks.

Other possible benefits of integrating PROs and DBs in AI-based models include providing personalized care and anticipating healthcare needs (individual benefit), cost reduction, and increased quality performance (public benefit) [42].

In summary, the burden of collecting both PROs and DBs must be weighed against the benefits of collecting them.

For AI adoption in healthcare to yield individual and public benefits, it is crucial to ensure seamless interoperability and effective data sharing across systems. Interoperability, the ability of different healthcare systems and technologies to communicate and exchange data, is essential for the efficient deployment and integration of AI solutions [43].

In line with the principle of beneficence, interoperability and secure data sharing facilitates better communication between healthcare providers, patients, and AI systems, allowing for a more coordinated approach to care that maximizes benefits and minimizes risks [44].

Principle of Non-Maleficence

The principle of nonmaleficence is related to avoiding the causation of harm [25].

Healthcare data are some of the most sensitive information about a person. Respecting a person's privacy is a fundamental ethical principle because it is directly linked with patient autonomy, personal identity, and well-being [45]

PROs and DBs, as healthcare-sensitive information, raise ethical, legal and social concerns regarding their collection, storage, access, and sharing [18].

Specifically, DBs data is collected and shared from personal devices, whose encryption methods may

not be sufficient to guarantee data privacy [46]. Therefore, data must be stored securely, preventing unauthorized access and possible misuse (also considering respect for the principle of non-maleficence). Several regulations [36, 47, 48] emphasize implementing procedures to guarantee data confidentiality. Furthermore, data controllers should carefully determine access levels granted to other stakeholders while considering and respecting autonomy, the right to privacy, and legitimate uses of data [34]. Also, individuals must maintain control over their personal data, especially regarding its collection, use, and dissemination [36, 47].

These and other concerns led to important initiatives such as the Global Pulse Data Privacy and Data Protection Principles from the United Nations and the General Data Protection Regulation (GDPR) from the European Union, both adopted in 2018 [31]. The GDPR establishes a legal foundation for data protection by encouraging a more proactive, systematic, and comprehensive approach [32]. This regulation thus expresses a legislative initiative to address emerging digital challenges regarding personal data [49]. The European Union (EU) Data Act, which entered into force on January 11, 2024, and complements the Data Governance Act [50], is also a significant development in the European Union's data landscape that aims to ensure fair access to and use of data across all economic sectors, including healthcare, while ensuring protection of personal data [48].

The Data Act aligns with the European Health Data Space (EHDS), providing secure access to a wide range of health data. Its main goal is to empower individuals to take ownership of their health data while allowing the EU to fully utilize the potential provided by a safe and secure exchange, use, and reuse of health data [51].

In addition, the AI Act, proposed by the European Commission in 2021 and unanimously endorsed by the EU 27 member states in December 2023 and recently entered into force, is the first comprehensive AI law to regulate AI use in the European Union [52]. This legal document addresses several ethical, legal and social implications of AI, including those related to adopting AI-based models [52]. The AI Act classifies AI systems based on a 'risk-based approach' and mandates various development and use requirements in multiple sectors, including healthcare [53]. It emphasizes the need for transparency in AI systems and addresses privacy and security concerns [53], which would be relevant to data preprocessing.

Even with ongoing improvements in AI and data accessibility, AI-based solutions in healthcare may be linked to malfunctions and errors that raise questions about patient safety, including (1) false negatives, such as missed diagnoses of life-threatening conditions; (2) false positives, leading to unnecessary treatments; and (3) inappropriate interventions, for example, caused by inaccurate

diagnosis or incorrect prioritization of interventions [54]. These situations might be related to differences between data used to train an algorithm and real-world data encountered during clinical deployment, highlighting the need for broad testing and validating AI algorithms in diverse patient populations [30].

When collecting DBs for AI-based models predicting patients' outcomes and taking actions on such basis, the risk of misdiagnosis must be considered and mitigated. Wearable devices can be inappropriately calibrated or malfunction, resulting in wrong readouts with a potentially significant impact on individuals and healthcare systems, contributing to increased stress and anxiety, waste of resources, and reduced quality of health provision [4]. To minimize the risks associated with the use of data originating from wearables, accreditation standards should be developed [49]. In the European Union, there is still no separate regulation for wearables, so these devices can be classified and certified as medical devices in accordance with Regulation (EU) 2017/745 [55]. This regulation establishes the rules for market introduction, availability, and entry into service of medical devices for human use and their accessories. Medical devices are classified according to their intended purpose and their intrinsic risks [55].

Principle of Justice

The ethical principle of justice is intimately related to fairly distributing benefits, risks, and costs [25, 56]. In this context, ethical considerations concern equitable access to AI solutions. Digital health tools might not be accessible to vulnerable communities and populations [57]. Furthermore, the design and development of digital health products and services might not adequately consider the needs of these populations [34]. Specifically, DBs are collected through wearables that might not be accessible to everyone, reinforcing inequalities [32]. Thus, ensuring fairness and equity in access is critical since the planning/design phase of digital health tools [34].

The fact that AI-based models are trained through large datasets might exacerbate embedded systematic health and social biases [58]. The underrepresentation of certain groups in training and testing datasets used to create and validate AI models is another factor contributing to bias in AI [30]. Therefore, the design of AI algorithms based on PROs and DBs should consider the need for data debiasing and contextualization to avoid discrimination and stigmatization of more vulnerable groups [59].

Training AI algorithms with real-world data and high-quality scientific evidence data might reduce bias and improve clinical practices [41].

Accordingly, the Montréal Declaration [39] states that AI systems must be designed and trained to create a just and equitable society, ensuring that it does not create, reinforce, or reproduce discrimination. Furthermore, other regulatory efforts refer to the same recommendation [23, 60–62].

These concerns might be addressed from the beginning of the design process, for instance, by ensuring a multidisciplinary and participatory approach [63].

Principles of Transparency and Accountability

Respect for the principles of transparency and accountability should also be highlighted in the context of AI adoption in healthcare [26].

Transparency consists of opening the entire process of AI adoption to public scrutiny, including decision-making and developed actions [20]. Accountability refers to the obligation to explain and justify the actions taken [64] and to establish mechanisms to ensure that responsibilities are traced and held [34]. These two principles are especially critical when analyzing ethical, legal, and social issues regarding AI processes [54].

AI transparency is closely linked to traceability and explainability/interpretability [54]. Traceability refers to transparently documenting the whole AI development process, including monitoring the model's performance in real-world practice after deployment [54]. Meanwhile, explainability or interpretability refers to the capacity to comprehend and convey AI systems' reasoning, choices, and actions. This is especially crucial in the healthcare industry, as AI-powered systems have the potential to impact critical decisions like prognosis and diagnosis [65]

Low interpretability is usually related to more complex models, such as neural networks, where comprehending the relation between the input and the output is challenging and may raise concerns [66–68]. For instance, a physician may know that a patient is at risk without fully understanding the underlying factors contributing to that risk and, consequently, how to address it effectively [38].

The “black box” nature of AI-based models may reflect a lack of transparency, even to developers [41, 45], which may result in public mistrust and social rejection of AI implementation at consequent opportunity costs [4]. In addition, with less interpretable models, identifying and correcting possible sources of bias can be more challenging [38, 45, 68]. Therefore, transparency regarding structures, underlying models, stakeholders, and their interests is a key value [34] that contributes to trustworthiness [26] and reinforces respect for other ethical principles [69].

Notwithstanding, AI-based models' continuous learning and adaption over time involves keeping up

with new data [38], which reflects changes in populations, technology, and care processes [41]. Therefore, subsequent evaluations are necessary even if an AI algorithm has been subjected to a rigorous initial evaluation [41].

The United States Food and Drug Administration (FDA) [70] presented innovative regulations to monitor modifications in adaptive models. These regulations require manufacturers to prespecify the anticipated modifications and establish protocols to address possible risks from these modifications [70]. Furthermore, the involvement of the public (including patients and providers) in the design of AI-based models may help mitigate these risks [4].

Accountability is also central to trustworthy and applicable AI in the healthcare field. Nonetheless, there are still legal gaps in the current national and international laws about who should be responsible or liable for errors or malfunctions of AI systems [54]. Since there are many actors engaged in the development and implementation of AI solutions for healthcare, such as patients, healthcare professionals, developers, and healthcare organizations, it is challenging to define roles and responsibilities [4, 54]. Structured frameworks and mechanisms are essential for appropriately assigning responsibility to all participants involved in the healthcare-AI workflow, incentivizing the implementation of comprehensive measures and best practices to mitigate errors and reduce harm to patients [71]. Another recommended approach to reinforce accountability could consist of regular audits and risk assessments to determine the level of regulatory oversight required for a specific AI tool [54]. These assessments should encompass the entire AI pipeline, from data collection and development to pre-clinical stages, deployment, and ongoing use [54].

Phase 2: Modified Delphi Study

The modified Delphi study included 23 experts in the first round and 25 experts in the second and final round, from diverse professional backgrounds and demographics. Table 2 summarizes the final round participants' characteristics.

Table 2. Participant Characteristics (Final Round, n = 25)

Category	Subgroup	n	%
Professional Background	Professors/Researchers	9	36.0%
	Technology Company Members	8	32.0%
	Healthcare Professionals	3	12.0%
	Data Protection Officers/Regulators	3	12.0%
	Health Administrators	2	8.0%
Gender	Male	16	64.0%

	Female	9	36.0%
Age Group	18-30	2	8.0%
	31-40	5	20.0%
	41-50	9	36.0%
	51-60	5	20.0%
	61-70	2	8.0%
	71+	2	8.0%
Country	Portugal	20	80.0%
	United Kingdom	1	4.0%
	Croatia	1	4.0%
	Italy	1	4.0%
	Lithuania	1	4.0%
	Spain	1	4.0%

In the final round, participants rated 64 items (as shown in Table 1). Of these, 57 items (89.1%) reached a high level of consensus, defined as $\geq 80\%$ agreement, with a median rating of 5 ("strongly agree"). The 7 items that did not reach consensus are the following:

- Periodically review and reaffirm consent mechanisms to ensure sustained autonomy in long-term data collection. (Agreement:72%)
- Inform individuals about how their data was used in AI model development, including details on what data was included or excluded. (Agreement:76%)
- Establish mechanisms that allow patients to contest AI-generated insights or decisions they disagree with. (Agreement:72%)
- Enable patients to revise their inputs over time, ensuring AI models adapt without introducing bias or undue influence. (Agreement:72%)
- Use AI to support patient advocacy by empowering individuals to analyze their own data, identify trends, and generate reports for healthcare consultations. (Agreement: 68%)
- Regularly publish publicly accessible reports on AI system performance to build trust and demonstrate continuous improvement. (Agreement: 64%)
- Make information about the databases and statistical methodologies used in AI training and validation publicly available. (Agreement: 76%)

Table 3 presents the recommendations that reached a high level of consensus ($\geq 80\%$) along with their median ratings and interquartile ranges (IQR). These recommendations, grouped according to core ethical principles, reflect the expert consensus on the essential aspects of integrating PROs and DBs in AI healthcare solutions. The table highlights key considerations for autonomy, beneficence, non-maleficence, justice, transparency, and accountability, providing actionable guidance for ethical and effective implementation.

Preprint
JMIR Publications

Table 3. Key Recommendations for Integrating PROs and DBs in AI Healthcare Solutions

Principle	Recommendation	Level of Agreement (%)	Median (IQR)
Autonomy	Ensure consent processes are comprehensible and adaptable to different levels of health literacy, promoting equitable patient engagement.	96%	5 (5-5)
	Adopt dynamic and continuously updated consent models, ensuring clear and tailored communication about data collection, usage, and implications.	92%	5 (4-5)
	Establish transparent protocols for patient communication when AI-driven decisions impact their health, ensuring information is clear, accurate, and precise.	92%	5 (4-5)
	Grant patients the "right to explanation," enabling them to understand the rationale behind AI-generated recommendations.	88%	5 (4-5)
	Share information about the algorithm's performance, the measures taken to identify and prevent errors, and the repercussions for patients' health if the AI algorithm is biased or incorrect.	88%	4 (4-5)
	Provide individuals control over their data, including the ability to view, manage, modify, and withdraw consent at any time.	84%	5 (4-5)
	Standardize consent processes across national and European levels to ensure uniformity and ease of patient authorization.	84%	5 (4-5)
	Implement granular consent mechanisms, allowing individuals to authorize specific uses of their data.	80%	5 (4-5)
Beneficence	Validate AI model effectiveness across different patient groups and stakeholders at multiple checkpoints over time, rather than relying solely on initial assessments.	100%	5 (4-5)
	Establish protocols for ongoing assessment and validation of ML models to ensure they continue to deliver beneficial outcomes.	100%	5 (4-5)
	Support responsible AI development by implementing audits and certifications without stifling innovation. Provide mechanisms for developers to promote responsible AI development,	100%	5 (4-5)

	reducing barriers without compromising safety or ethical standards.		
	Implement a comprehensive evaluation framework that tracks performance at each stage of the AI pipeline (data acquisition, preprocessing, modeling, and deployment) to identify sources of deviation and necessary adjustments.	96%	5 (4-5)
	Uphold human oversight as a core principle, ensuring AI-driven healthcare decisions are reviewed by qualified professionals.	92%	5 (4-5)
	Ensure interoperability between PROs and DBs platforms with other healthcare systems, incorporating real-time data processing capabilities to facilitate timely and informed decision-making for both clinicians and patients.	88%	5 (4-5)
	Facilitate secure and ethical patient data sharing across organizations to enhance accuracy and reliability in patient care while safeguarding privacy.	88%	5 (4-5)
	Implement continuous model and data drift monitoring with predefined thresholds to trigger model retraining or resampling when necessary.	88%	4 (4-5)
	Regularly update and recalibrate models to reflect changes in patient populations and healthcare practices.	88%	5 (4-5)
	Assess AI system benefits from multiple stakeholder perspectives, recognizing distinct priorities among patients, healthcare providers, and administrators.	80%	4 (4-5)
	Safeguard human security through robust cybersecurity measures, preventing unauthorized data access and mitigating potential misuse.	100%	5 (5-5)
Non-maleficence	Train staff on cybersecurity best practices, ensuring that individuals involved in the AI lifecycle understand their responsibility to protect patient data from unauthorized access and misuse.	100%	5 (5-5)
	Establish protocols to address data breaches, including audits to determine the causes and prevent recurrence, and implement more secure procedures to protect patient data.	100%	5 (4-5)
	Develop alert mechanisms that flag potential errors or inconsistencies in AI	100%	5 (4-5)

	outputs, allowing for timely interventions to prevent harm.		
	Conduct regular impact assessments to evaluate the effects of AI solutions, including patient outcomes (e.g., satisfaction, health improvements), process efficiency, and potential negative consequences (e.g., patient harm, errors).	100%	5 (5-5)
	Apply "security by design" principles, throughout the AI development process, rather than relying solely on policies, to ensure that security measures are inherently built into the technology.	96%	5 (4-5)
	Integrate privacy and data sovereignty measures into AI solutions, ensuring that patient privacy is maintained without compromising scalability or operational efficiency.	96%	5 (4-5)
	Implement stringent safeguards to protect stored patient data from unauthorized access.	96%	5 (5-5)
	Perform data and model audits using synthetic data to identify potential vulnerabilities, including security issues or model biases, and ensure that ethical and legal standards are maintained.	96%	5 (4-5)
	Implement comprehensive audits and oversight systems to ensure ongoing compliance with data protection regulations (e.g., GDPR) and the AI Act, and assess the effectiveness of decisions made based on AI outputs.	92%	5 (4-5)
	Ensure AI applications in healthcare undergo Data Protection Impact Assessments (DPIAs) as required by GDPR Article 35.	92%	5 (4-5)
	Ensure AI systems provide explainable outputs and maintain human oversight in all healthcare decisions to prevent reliance on AI as a sole decision-maker. Ethical skepticism should guide the integration of AI into decision-making processes.	88%	5 (4-5)
	Establish clear notification protocols in case of security breaches or data leaks, ensuring patients are promptly informed.	80%	4 (4-5)

Justice	Maintain transparency and accountability in AI systems so patients understand how their data is used and trust that it contributes to fair, unbiased outcomes.	96%	5 (4-5)
	Actively involve diverse populations, including those from varying cultural, socioeconomic, and medical backgrounds, in the design, development, and implementation of digital health tools to ensure equitable access and outcomes.	92%	5 (4-5)
	Ensure that digital health tools are accessible to all populations, particularly vulnerable and underserved communities, with particular attention to mitigating the underuse of these tools by these groups.	92%	5 (4,8-5)
	Ensure training datasets undergo rigorous bias and representativity assessments to achieve fair outcomes across all populations.	92%	5 (4-5)
	Ensure that underrepresented populations are not harmed by biases in ML models, and that measures are in place to regularly detect and correct bias during model development and deployment.	92%	5 (4-5)
	Avoid any form of positive or negative discrimination in feature selection and model design by ensuring that all categories or features used in AI models are balanced, representative, and free from bias.	92%	5 (4-5)
	Implement legal frameworks and harmonize regulations across the EU and member states to create an environment where AI development is both ethical and just, with a focus on ensuring equal access to AI benefits.	92%	5 (4-5)
	Balance AI regulation to protect patients without stifling innovation, fostering an environment that supports both safety and competitiveness in the EU AI ecosystem.	92%	5 (5-5)
	Ensure diverse stakeholder engagement in designing and implementing ML models to mitigate biases.	88%	4 (4-5)
	Monitor and analyze the demographic profile of populations using digital health tools and take proactive steps to mitigate barriers to access for vulnerable and	80%	5 (4-5)

	underserved groups.		
Transparency and Accountability	Ensure human oversight is incorporated into AI systems to guarantee accountability and maintain a role for healthcare providers in decision-making.	100%	5 (4-5)
	Establish regular and well-defined audit intervals for AI systems that ensure continuous compliance with ethical and legal standards, while maintaining a balanced approach that upholds accountability without hindering innovation in AI development.	96%	5 (4-5)
	Conduct regular audits and risk assessments to identify and mitigate potential errors and malfunctions.	96%	4 (4-5)
	Implement and regularly update training programs for users of AI healthcare tools to ensure they understand the principles of transparency and accountability.	96%	5 (4-5)
	Ensure transparency in AI integration, informing both patients and healthcare professionals of its role and impact in decision-making.	96%	5 (4-5)
	Implement measures to ensure that AI development and deployment processes, including decision-making criteria and model performance, are well-documented and openly shared with relevant stakeholders.	92%	5 (4-5)
	Document and openly share the AI development and deployment process, including decision-making criteria and model performance.	92%	5 (4-5)
	Encourage transparent reporting of AI-related errors and unintended consequences by both healthcare providers and AI developers.	92%	5 (4-5)
	Integrate MLOps best practices in AI-driven healthcare solutions to ensure robust, efficient, and transparent model development, deployment, and maintenance.	92%	4 (4-5)
	Ensure AI outputs support, rather than dictate, healthcare decisions.	92%	5 (4-5)
	Clearly define the responsibility of regulatory bodies and oversight entities in the continuous evaluation and governance of AI models in healthcare.	92%	5 (4-5)
	Ensure that all stakeholders, including patients and healthcare providers, are	88%	5 (4-5)

	informed about any errors or malfunctions that may negatively impact patient care.		
	Ensure AI models provide interpretable outputs, including the probability and uncertainty of results, to enhance user understanding.	88%	5 (4-5)
	Enhance the interpretability of ML models to ensure healthcare providers and patients can understand and trust their outputs.	88%	5 (4-5)
	Establish data lineage protocols to track all data actions, ensuring a non-repudiation method for authorized actors involved in AI healthcare workflows.	88%	4 (4-5)
	Promptly disclose any issues related to patient data usage, such as breaches or misuse, to maintain trust in AI applications.	84%	5 (4-5)

Figure 2 provides a visual summary of the main experts' recommendations for the integration of PROs and DBs in AI-based models and their respective correspondence to the main ethical principles as resulting from our narrative literature review.



Figure 2. Experts' recommendations for integration of PROs and DBs in AI-based models and their respective correspondence to the main ethical principles as resulting from our narrative literature review. [Author's elaboration]

DISCUSSION

This study presents the first international expert consensus specifically addressing the ethical integration of PROs and DBs in AI healthcare models, exploring its ethical, legal, and social considerations. While existing literature extensively covers general AI ethics principles, these findings reveal unique considerations that arise specifically from patient-generated health data. The 57 consensus recommendations provide targeted guidance for the distinct challenges posed by continuous patient data collection.

The findings highlight a strong consensus among participants regarding the importance of upholding key ethical principles in this context. The healthcare landscape is indeed undergoing a significant transformation with the increasing integration of AI across various domains, further fueled by the growing availability and utilization of PROs and DBs. PROs capture patients' perspectives on their health status, symptoms, and functional abilities, offering valuable subjective insights, while DBs, from wearable sensors and other technologies, provides continuous and objective data. The convergence of AI with these rich data sources holds immense potential to revolutionize healthcare by enabling personalized interventions, predictive analytics, and more efficient clinical workflows. However, this powerful integration also brings forth a complex array of ethical, legal, and social implications (ELSI) that demand careful

consideration to ensure responsible innovation and safeguard patient rights [72].

Principle of Autonomy

Regarding the principle of autonomy, participants emphasized the need for individuals to maintain control over their data, including the ability to view, manage, modify, and withdraw consent at any time. This aligns with the ethical principle of respecting and supporting autonomous decisions [25]. The importance of clear and tailored communication, along with ensuring understandable consent processes that accommodate varying levels of health literacy, was strongly affirmed. This is consistent with the notion that AI adoption should contribute to patient-centered healthcare, where patients are involved in decision-making processes and are well-informed about the purpose, risks, benefits, and alternatives of any intervention. Research consistently highlights informed consent and patient autonomy as paramount ethical considerations in this rapidly evolving field [24]. Patients must be informed about the role of AI in their care by healthcare professionals and technology developers and should have the freedom to decline its involvement if they feel uncomfortable. This perspective aligns with the understanding that AI should serve as a tool to augment, rather than replace, the essential human interaction and empathy that are integral to quality patient care [73, 74].

Several recommendations focused on enhancing the informed consent process. The call for standardized consent across national and European levels reflects the need for uniformity and ease of patient authorization. Furthermore, the emphasis on dynamic and continuously updated consent models, along with transparent protocols for patient communication, underscores the importance of ensuring continued comprehension and voluntariness. The recognition that informed consent comprehension can be influenced by factors such as age, gender, health literacy, and contextual factors [75] highlights the need for adaptable consent processes. This is particularly crucial for collecting DBs, where continuous passive data collection can potentially limit autonomy due to unanticipated data treatment.

The concept of a “right to explanation,” enabling patients to understand the reasoning behind AI system recommendations, received substantial support (88%). This aligns with the broader principle of transparency, which is crucial for building trust in AI-based healthcare models[76]. Sharing information about the algorithm's performance,

error prevention measures, and potential repercussions for patients' health in cases of bias or errors was also considered important. This "right to explanation" has gained significant traction, emphasizing the need for explainable AI (XAI) to foster patient trust and informed decisions [72, 77].

Principle of Beneficence

The principle of beneficence focused on promoting well-being and balancing benefits against risks and costs, also garnered strong agreement. Participants strongly supported the implementation of a comprehensive evaluation framework to track AI pipeline performance (96%) and the establishment of protocols for ongoing assessment and validation of machine learning models (100%). This reflects the understanding that the continuous learning characteristic of AI algorithms necessitates continuous evaluation to ensure ongoing patient benefit [78].

Recommendations emphasized the importance of validating AI model effectiveness across diverse patient groups and stakeholders over time rather than relying solely on initial assessments (100%). The need for continuous model and data drift monitoring, along with regular updates and recalibration of models, was also highlighted (88%). Ensuring interoperability between PROs and DBs platforms with other healthcare systems and promoting secure and ethical data sharing across organizations were identified as key to improving the accuracy and reliability of patient care (88%). These measures align with the potential benefits of integrating PROs and DBs in AI-based models, such as providing personalized care, anticipating healthcare needs (individual benefit), cost reduction, and increased quality performance (public benefit). Robust frameworks for evaluating AI pipeline performance are essential to ensure beneficence [79]. Continuous assessment is crucial due to AI's learning nature [80].

The recommendation for human oversight was strongly endorsed (92%), reinforcing the importance of qualified professionals reviewing AI-driven decisions in healthcare. Furthermore, participants acknowledged the need to balance audits and certifications with support mechanisms for developers to promote responsible AI development without stifling innovation (100%). While AI offers significant benefits, human oversight remains critical to ensure patient safety and maintain the human aspects of care, like empathy [24]. Patients may find it difficult to accept fully automated "machine-human" medical relationships [24].

Principle of Non-Maleficence

The principle of non-maleficence, centered on avoiding harm, was addressed through recommendations focused on data privacy and security, monitoring burden, and safety. Participants strongly agreed on the need to conduct regular impact assessments to evaluate the effects of AI solutions (100%) and to design AI systems that ensure explainability of their outputs (88%). The importance of human oversight in all healthcare decisions, with AI serving as a support tool rather than a sole decision-maker, was emphasized (88%). This aligns with the recognition that AI-based solutions may be linked to malfunctions and errors that raise questions about patient safety [81].

To mitigate potential harm, participants recommended developing alert mechanisms to flag errors or inconsistencies in AI outputs (100%) and establishing protocols to address data breaches, including audits and more secure procedures to protect patient data (100%). The need to safeguard human security alongside technological measures, prevent unauthorized access to data, and invest in comprehensive cybersecurity practices was strongly affirmed (100%). Training staff on cybersecurity best practices and implementing comprehensive audits and oversight systems to ensure ongoing compliance with data protection regulations were also deemed crucial (100%, 92%). Robust security protocols and cybersecurity measures are essential to prevent data breaches [82].

The principles of "security by design" and "privacy and data sovereignty" should be integrated into AI solutions (96%). Furthermore, AI applications in healthcare should be subject to Data Protection Impact Assessments (DPIAs) as per Article 35 of the GDPR (92%). The security of stored patient data must be explicitly addressed with stringent safeguards (96%). Data and model audits using synthetic data were recommended to identify potential vulnerabilities and ensure adherence to ethical and legal standards (96%).

Federated learning emerges as a promising approach that supports both beneficence and non-maleficence principles. Allowing AI models to be trained across multiple decentralized sites without exchanging the raw data itself, federated learning enables collaborative model development while preserving data privacy [83]. Participants recognized that secure and ethical patient data sharing should be facilitated across

organizations to enhance accuracy and reliability in patient care while safeguarding privacy (88%). This technique particularly benefits healthcare organizations seeking to develop robust AI models by leveraging diverse patient populations across institutions while maintaining strict data governance and sovereignty [84], which was reinforced by the recommendation to integrate privacy and data sovereignty measures into AI solutions, ensuring that patient privacy is maintained without compromising scalability or operational efficiency (96%).

These recommendations reflect the sensitive nature of healthcare data and the ethical concerns regarding its collection, storage, access, and sharing [72]. The lack of transparency in some AI systems ("black box" phenomenon) also poses a risk of harm [85]. Explainability is crucial for identifying errors and ensuring correct decisions [85–87]

Principle of Justice

The principle of justice, focused on the fair distribution of benefits, risks, and costs, was addressed through recommendations aimed at ensuring equitable access to AI solutions and mitigating biases. Participants emphasized the importance of diverse stakeholder engagement in designing and implementing machine learning models (88%) and actively involving diverse populations in the design, development, and implementation of digital health tools (92%). This aligns with the need to consider the needs of vulnerable communities and populations and to ensure fairness and equity from the planning/design phase [88].

Training datasets should undergo rigorous bias and representativity assessments (92%), and measures should be in place to regularly detect and correct bias during model development and deployment (92%). This is crucial to avoid discrimination and stigmatization of vulnerable groups, as AI-based models trained on large datasets may exacerbate embedded systematic health and social biases. Ensuring that digital health tools are accessible to all populations, particularly vulnerable and underserved communities, was also considered essential (92%). Identifying and mitigating biases in AI algorithms and training datasets is vital for justice and equity [89].

Transparency and accountability in AI systems were deemed necessary to allow patients to understand how their data is used and to ensure fair and unbiased outcomes (96%).

Participants also recommended avoiding any form of positive or negative discrimination in feature selection and model design (92%) and implementing legal frameworks and harmonizing regulations across the EU and member states to create an environment where AI development is both ethical and just (92%). Balancing the enforcement of EU regulations with the competitiveness and ease of use for AI industry stakeholders in Europe was identified as a key challenge (92%). Diverse stakeholder engagement is crucial for achieving fair outcomes [77].

Principles of Transparency and Accountability

The principles of transparency and accountability were consistently highlighted across recommendations. There was strong consensus on the need for interpretable AI outputs (88%), human oversight in AI systems (88-100%), and well-documented development processes (92%). Experts agreed on the importance of establishing data lineage protocols (88%), implementing training programs for AI users (96%), and clearly defining regulatory responsibilities (92%).

These recommendations strongly align with established frameworks in responsible AI governance. The emphasis on interpretability addresses the "black box problem" in healthcare AI, where low interpretability raises ethical concerns around trust and clinical adoption [90].

Also, the high consensus around error reporting (92%) and stakeholder communication about any errors or malfunctions that may negatively impact patient care (88%) aligns with work by Amann et al. [87], emphasizing that explainability is not merely a technical challenge but a multidisciplinary requirement involving clinical, technical, legal, and patient perspectives.

The recommendations for regular audits and transparent documentation align with the need for traceability, explainability/interpretability, and mechanisms to ensure that responsibilities are traced and held in the development and implementation of AI solutions for healthcare [91].

In addition to the principles discussed below, our findings implicitly address the need for adherence to FAIR data principles (Findability, Accessibility, Interoperability, and Reusability) [92].

Findability is supported through recommendations for establishing data lineage protocols (88% agreement) and documenting AI development processes (92%

agreement). These practices ensure that data can be discovered and identified through proper metadata and persistent identifiers, which are critical for tracing data provenance in AI systems [93].

Accessibility is reflected in recommendations for transparent data sharing across organizations (88% agreement) and ensuring that data access is secured through appropriate authentication and authorization protocols. The emphasis on protocols to address data breaches (100% agreement) and implementation of comprehensive audits (92% agreement) support controlled accessibility that balances openness with security.

Interoperability received explicit support (88% agreement) through recommendations to ensure PROs and DBs platforms can communicate with other healthcare systems. This facilitates the exchange and integration of data across different platforms and applications, enhancing the potential for comprehensive analysis while maintaining semantic integrity [94].

Reusability is embedded in recommendations for comprehensive documentation of AI development processes (92% agreement) and establishing protocols for ongoing assessment and validation (100% agreement).

The integration of FAIR principles with ethical considerations creates a robust framework for responsible AI development in healthcare, supporting both innovation and patient protection [95].

Beyond the core ethical principles analyzed in our study, there are complementary considerations that merit attention in future research on integrating PROs and DBs in AI healthcare solutions. One such consideration relates to the data collection burden itself.

The potential burden on individuals responding to PRO questionnaires [96], represents an important non-maleficence consideration that deserves further investigation. Factors such as the length or formatting of the questionnaire, the literacy level of respondents, and the mode of administration are potential factors related to response burden [97]. Similarly, for DBs, there may be associated burdens related to carrying or safeguarding electronic devices and the discomfort some patients experience from being tracked [98].

Additionally, tight follow-up protocols and subsequent interventions may influence patient responses, potentially compromising data validity and the benefits of AI algorithms. For instance, individuals may deliberately omit a response to avoid contact or specific clinical interventions [38]. These considerations highlight the critical

importance of involving healthcare professionals in designing tailored interventions that respect individual patient preferences and behaviors, ensuring inclusive and ethical implementation [99].

These dimensions of data collection burden provide relevant context to our findings on ethical integration of PROs and DBs in AI healthcare solutions and could be further elaborated in future research.

Strengths and Limitations

Although several studies address the ethical considerations in the use of artificial intelligence in healthcare [21, 100–104], this is, to the best of our knowledge, the first study to comprehensively map and to achieve international expert consensus specifically on PROs and DBs integration in AI healthcare models. This targeted approach distinguishes our work from broader AI ethics frameworks and provides actionable guidance for the specific challenges that developers, researchers, and policymakers face when integrating PROs and DBs in AI-based healthcare models. This novel contribution enhances the relevance of the findings beyond this specific context. While the analysis is centered on PROs and DBs, the identified considerations are relevant to a broad range of health-related data processing scenarios. As such, they can serve as a foundational ethical and operational framework for policymakers, developers, researchers, and practitioners working across different health technologies.

Methodologically, the use of a Delphi survey provided a structured way to build expert consensus on complex normative issues [105]. The process benefited from the diversity of participants across disciplines and geographies, which enriched the discussion and supported a pluralistic view of the challenges and priorities involved. Importantly, the survey was designed to ensure full anonymity; no login credentials or personal information that could allow the identification of the participants were collected at any point. While this approach safeguarded privacy and encouraged open contributions, it also introduced certain limitations. Specifically, the research team could not track individual participation across rounds, limiting the ability to assess response consistency, drop-out rates, or potential shifts in opinion.

Furthermore, the study does not include the perspectives of patients or the wider public, which may have provided additional insights, particularly around how data-driven

healthcare models are perceived by those who are one of the most important beneficiaries [106]. Including these perspectives in future research will be essential for a more comprehensive understanding of public trust, expectations, and concerns.

Finally, the findings reflect a predominantly European ethical and regulatory context. While many of the issues raised have broad relevance, their implications and proposed solutions may differ in settings with distinct legal systems, healthcare structures, or sociocultural norms [107].

CONCLUSION

The large-scale processing of patient-generated health data, particularly PROs and DBs, by ML algorithms presents unique ethical challenges. This study provides the first international consensus framework specifically designed for these data types, analyzing ethical considerations through traditional bioethical principles complemented by transparency and accountability.

In this study, we have analyzed the ethical, legal, and social considerations regarding integrating PROs and DBs in ML models based on the four traditional bioethical principles of respect for autonomy, beneficence, non-maleficence, and justice, complemented by the principles of transparency and accountability.

Despite our analysis being focused on PROs and DBs, the identified considerations can be used as a baseline for a wide range of contexts of health-related data processing, supporting policymakers, developers, researchers, and practitioners in tackling these potential risks.

It is crucial to consider the ethical threats and risks associated with integrating sensitive personal information into ML models, explore its potential impact on every aspect of people's lives, and build regulation systems that enhance and protect their rights.

This study offers a timely and relevant contribution to ongoing efforts to align technological innovation with ethical and societal values. It provides a structured, consensus-driven foundation that can inform both policy development and practical implementation of AI in health systems that increasingly rely on diverse and dynamic data sources.

DECLARATIONS

Abbreviations

AI: Artificial Intelligence; GDPR: General Data Protection Regulation; PROs: Patient-reported outcomes; DBs: Digital Biomarkers.

Ethics approval and consent to participate

Not applicable.

Consent for publication

Not applicable

Competing Interests

None.

Funding

This research was funded by Fundação para a Ciência e Tecnologia (FCT, Portugal) through national funds to the Associated Laboratory in Translation and Innovation Towards Global Health REAL (LA/P/0117/2020).

Authors' contributions

JS conceived the idea, conducted the research, and wrote the original draft of the manuscript. JVC critically reviewed the manuscript and recommended additional references, significantly improving the manuscript. RS and TM reviewed the manuscript. All authors have approved the final manuscript.

Acknowledgements

The authors are grateful to all the participants who contributed their time and expertise to this study.

REFERENCES

1. Callahan A, Shah NH. Machine Learning in Healthcare. In: Key Advances in Clinical Informatics: Transforming Health Care Through Health Information Technology. Elsevier Inc.; 2017. p. 279–91.
2. Uttiramerur A. The Role of Machine Learning in Identifying Risk Factors for Adverse Events in Healthcare: A Comprehensive Overview Programmer Analyst at

- ThermoFisher Scientific, USA. Arvind Uttiramerur, Programmer Analyst at ThermoFisher Scientific. 2023;2:1–3.
3. Drabiak K, Kyzer S, Nemov V, El Naqa I. AI and machine learning ethics, law, diversity, and global impact. *British Journal of Radiology*. 2023;96.
 4. Morley J, Machado CCV, Burr C, Cowls J, Joshi I, Taddeo M, et al. The ethics of AI in health care: a mapping review. *Soc Sci Med*. 2020;260 June.
 5. Familoni BT. Ethical frameworks for AI in healthcare entrepreneurship: A theoretical examination of challenges and approaches. *International Journal of Frontiers in Biology and Pharmacy Research*. 2024;5:057–65.
 6. Browne JP, Cano SJ, Smith S. Using Patient-reported Outcome Measures to Improve Health Care: Time for a New Approach. *Med Care*. 2017;55:901–4.
 7. Aiyegbusi OL, Hughes SE, Calvert MJ. The Role of Patient-Reported Outcomes (PROs) in the Improvement of Healthcare Delivery and Service. *Handbook of Quality of Life in Cancer*. 2022;:339–52.
 8. Field J, Holmes MM, Newell D. <p>PROMs data: can it be used to make decisions for individual patients? A narrative review</p>. *Patient Relat Outcome Meas*. 2019;Volume 10:233–41.
 9. Dias AG, Roberts CJ, Lippa J, Arora J, Lundström M, Rolfson O, et al. Benchmarking Outcomes That Matter Most to Patients: The Globe Programme. *EMJ Innovations* 22 2017. 2017;2:42–9.
 10. Menditto VG, Maraldo A, Barbadoro P, Maccaroni R, Salvi A, D’Errico MM, et al. Patient-Reported Outcome Measurements (PROMs) After Discharge From the Emergency Department: A Cross-Sectional Study. *J Patient Exp*. 2021;8.
 11. Franklin PD, Chenok KE, Lavalee D, Love R, Paxton L, Segal C, et al. Framework To Guide The Collection And Use Of Patient-Reported Outcome Measures In The Learning Healthcare System. *eGEMs (Generating Evidence & Methods to improve patient outcomes)*. 2017;5:17.
 12. Cohen AB, Mathews SC. The Digital Outcome Measure. *Digit Biomark*. 2018;2:94–105.
 13. Horn ME, Reinke EK, Mather RC, O’Donnell JD, George SZ. Electronic health record–integrated approach for collection of patient-reported outcome measures: a retrospective evaluation. *BMC Health Serv Res*. 2021;21:1–11.
 14. Izmailova ES, Wagner JA, Bakker JP, Kilian R, Ellis R, Ohri N. A proposed multi-domain, digital model for capturing functional status and health-related quality of life in oncology. *Clin Transl Sci*. 2024;17:e13712.
 15. Andreolettiid M, Haller L, Vayena E, Blasimmeid A. Mapping the ethical landscape of digital biomarkers: A scoping review. *PLOS Digital Health*. 2024;3:e0000519.
 16. Motahari-Nezhad H, Fgaier M, Abid MM, Péntek M, Gulácsi L, Zrubka Z. Digital Biomarker–Based Studies: Scoping Review of Systematic Reviews. *JMIR Mhealth Uhealth*. 2022;10.
 17. Powell D. Walk, talk, think, see and feel: harnessing the power of digital biomarkers in healthcare. *NPJ Digit Med*. 2024;7.
 18. Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC. *Official Journal of the European Union*. 2016.
 19. European Health Data Space - European Commission. https://health.ec.europa.eu/ehealth-digital-health-and-care/european-health-data-space_en. Accessed 19 Jun 2024.
 20. Morley J, Machado CCV, Burr C, Cowls J, Joshi I, Taddeo M, et al. The ethics of AI

in health care: a mapping review. *Soc Sci Med*. 2020;260 June.

21. Murphy K, Di Ruggiero E, Upshur R, Willison DJ, Malhotra N, Cai JC, et al. Artificial intelligence for good health: a scoping review of the ethics literature. *BMC Med Ethics*. 2021;22:1–17.

22. World Health Organization. Ethics and Governance of Artificial Intelligence for Health: WHO guidance. Geneva: World Health Organization; 2021.

23. European Commission. Directorate-General for Communications Networks. Content and Technology. Ethics Guidelines for Trustworthy AI. Brussels; 2019.

24. Farhud DD, Zokaei S. Ethical Issues of Artificial Intelligence in Medicine and Healthcare. *Iran J Public Health*. 2021;50:i.

25. Beauchamp TL, Childress JF. Principles of Biomedical Ethics - Oxford University Press. 2019;:512.

26. Xafis V, Schaefer GO, Labude MK, Brassington I, Ballantyne A, Lim HY, et al. An Ethics Framework for Big Data in Health and Research. *Asian Bioeth Rev*. 2019;11:227–54.

27. Sharma T, Dhull A, Singh A, Singh KK. Designing transparent and accountable AI systems for healthcare. In: *Responsible and Explainable Artificial Intelligence in Healthcare: Ethics and Transparency at the Intersection*. Elsevier; 2024. p. 91–106.

28. Shang Z. Use of Delphi in health sciences research: A narrative review. *Medicine*. 2023;102:e32829.

29. World Medical Association. WMA Declaration of Helsinki – Ethical Principles for Medical Research Involving Human Participants – WMA – The World Medical Association. 2024.

30. Lehmann LS. Ethical Challenges of Integrating AI into Healthcare. *Artif Intell Med*. 2022;:139–44.

31. MacPherson Y, Pham K. Ethics in Health Data Science. In: Celi LA, Majumder MS, Ordóñez P, Osorio JS, Paik KE, Somai M, editors. *Leveraging Data Science for Global Health*. Springer; 2020.

32. Maher NA, Senders JT, Hulsbergen AFC, Lamba N, Parker M, Onnela J, et al. Passive data collection and use in healthcare: a systematic review of ethical issues. *Int J Med Inform*. 2019;129 May:242–7.

33. Obar JA, Oeldorf-Hirsch A. The Biggest Lie on the Internet: Ignoring the Privacy Policies and Terms of Service Policies of Social Networking Service. *Inf Commun Soc*. 2018;:1–20.

34. Brall C, Schröder-Bäck P, Maeckelberghe E. Ethical aspects of digital health from a justice point of view. *Eur J Public Health*. 2019;29:18–22.

35. Gesualdo F, Daverio M, Palazzani L, Dimitriou D, Diez-Domingo J, Fons-Martinez J, et al. Digital tools in the informed consent process: a systematic review. *BMC Med Ethics*. 2021;22:1–10.

36. European Union. Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation). 2016.

37. Art. 5 GDPR – Principles relating to processing of personal data - General Data Protection Regulation (GDPR). <https://gdpr-info.eu/art-5-gdpr/>. Accessed 22 Jun 2024.

38. Jacobson NC, Bentley KH, Walton A, Wang SB, Fortgang RG, Millner AJ, et al. Ethical dilemmas posed by mobile health and machine learning in psychiatry research. *Bull World Health Organ*. 2020;98:270–6.

39. Montréal declaration for a responsible development of artificial intelligence. 2018.

40. Beauchamp T. The Principle of Beneficence in Applied Ethics. Spring 2019 Edition.

41. Char DS, Abràmoff MD, Feudtner C. Identifying Ethical Considerations for Machine Learning Healthcare Applications. *American Journal of Bioethics*. 2020;20:7–17.
42. Cordeiro J V. Artificial Intelligence and Precision Public Health: A Balancing Act of Scientific Accuracy, Social Responsibility, and Community Engagement. *Portuguese Journal of Public Health*. 2024;42:1–5.
43. Joshi H. Enabling Next-Gen Healthcare: Advanced Interoperability and Integration with AI, IoMT, and Precision Medicine. *IARJSET*. 2024;8.
44. Torab-Miandoab A, Samad-Soltani T, Jodati A, Rezaei-Hachesu P. Interoperability of heterogeneous health information systems: a systematic literature review. *BMC Med Inform Decis Mak*. 2023;23:18.
45. Reddy S, Allan S, Coghlan S, Cooper P. A governance model for the application of AI in health care. *Journal of the American Medical Informatics Association*. 2020;27:491–7.
46. Segura Anaya LH, Alsadoon A, Costadopoulos N, Prasad PWC. Ethical Implications of User Perceptions of Wearable Devices. *Sci Eng Ethics*. 2018;24:1–28.
47. Montréal declaration for a responsible development of artificial intelligence. 2018.
48. European Commission. Regulation (EU) 2023/2854 of the European Parliament and of the Council of 13 December 2023 on harmonised rules on fair access to and use of data and amending Regulation (EU) 2017/2394 and Directive (EU) 2020/1828 (Data Act). 2023.
49. Vayena E, Blasimme A, Cohen IG. Machine learning in medicine: addressing ethical challenges. *PLoS Med*. 2018;15:4–7.
50. European Commission. Regulation (EU) 2022/868 of the European Parliament and of the Council of 30 May 2022 on European data governance and amending Regulation (EU) 2018/1724 (Data Governance Act). 2022.
51. European Commission. European Health Data Space. https://health.ec.europa.eu/ehealth-digital-health-and-care/european-health-data-space_en. Accessed 2 Apr 2024.
52. EU Artificial Intelligence Act | Up-to-date developments and analyses of the EU AI Act. 2024. <https://artificialintelligenceact.eu/>. Accessed 26 Mar 2024.
53. European Parliament. European Parliamentary Research Service. Artificial intelligence act. 2024.
54. European Parliament. European Parliamentary Research Service. Artificial intelligence in healthcare. Applications, risks and ethical and societal impacts. 2022.
55. European Commission. Regulation (EU) 2017/745 of the European Parliament and of the Council of 5 April 2017 on medical devices, amending Directive 2001/83/EC, Regulation (EC) No 178/2002 and Regulation (EC) No 1223/2009 and repealing Council Directives 90/385/EEC and 93/42/EEC. 2017.
56. Sen A. *The Idea of Justice*. The Belknap Press of Harvard University Press Cambridge, Massachusetts; 2009.
57. Brall C, Schröder-Bäck P, Maeckelberghe E. Ethical aspects of digital health from a justice point of view. *The European Journal of Public Health*. 2019;29 Suppl 3:18.
58. Murphy K, Di Ruggiero E, Upshur R, Willison DJ, Malhotra N, Cai JC, et al. Artificial intelligence for good health: a scoping review of the ethics literature. *BMC Med Ethics*. 2021;22:1–17.
59. Gasser U, Ienca M, Scheibner J, Sleight J, Vayena E. Digital tools against COVID-19: taxonomy, ethical challenges, and navigation aid. *Lancet Digit Health*. 2020;2:e425–34.
60. European Parliament. P9_TA(2024)0138 Artificial Intelligence Act European

Parliament legislative resolution of 13 March 2024 on the proposal for a regulation of the European Parliament and of the Council on laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act) and amending certain Union Legislative Acts (COM(2021)0206-C9-0146/2021-2021/0106(COD)) (Ordinary legislative procedure: first reading). 2024.

61. Executive Office of the President. Big Data: A Report on Algorithmic Systems, Opportunity, and Civil Rights. 2016.

62. The White House. Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence. 2023. <https://www.whitehouse.gov/briefing-room/presidential-actions/2023/10/30/executive-order-on-the-safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence/>. Accessed 21 Jun 2024.

63. Howard A, Borenstein J. The Ugly Truth About Ourselves and Our Robot Creations: The Problem of Bias and Social Inequity. *Sci Eng Ethics*. 2018;24:1521–36.

64. Bovens M. Analysing and assessing accountability: a conceptual framework. *European Law Journal*. 2007;13:447–68.

65. Li RC, Muthu N, Hernandez-Boussard T, Dash D, Shah NH. Explainability in Medical AI. 2022;:235–55.

66. Mckelvey TG, Ahmad MA, Eckert C, Teredesai A, Mckelvey G. Interpretable Machine Learning in Healthcare. 2018.

67. Zhang Y, Tiño P, Leonardis A, Tang K. A Survey on Neural Network Interpretability. 2021.

68. Markus AF, Kors JA, Rijnbeek PR. The role of explainability in creating trustworthy artificial intelligence for health care: A comprehensive survey of the terminology, design choices, and evaluation strategies. *J Biomed Inform*. 2021;113:103655.

69. Feudtner C, Schall T, Nathanson P, Berry J. Ethical framework for risk stratification and mitigation programs for children with medical complexity. *Pediatrics*. 2018;141 March 2018:S250–8.

70. US Food and Drug Administration (FDA). Artificial Intelligence/Machine Learning (AI/ML)-Based Software as a Medical Device (SaMD) Action Plan. International Medical Device Regulators Forum. 2021; January:1–9.

71. Shuaib A. Transforming Healthcare with AI: Promises, Pitfalls, and Pathways Forward. *Int J Gen Med*. 2024;17:1765–71.

72. Elendu C, Amaechi DC, Elendu TC, Jingwa KA, Okoye OK, Okah J, et al. Ethical implications of AI and robotics in healthcare: a review. *Medicine*. 2023. <https://doi.org/10.1097/MD.00000000000036671>.

73. Bajwa J, Munir U, Nori A, Williams B. Artificial intelligence in healthcare: transforming the practice of medicine. *Future Healthc J*. 2021;8:e188.

74. Dailah HG, Koriri M, Sabei A, Kriry T, Zakri M. Artificial Intelligence in Nursing: Technological Benefits to Nurse's Mental Health and Patient Care Quality. *Healthcare*. 2024;12:2555.

75. Gesualdo F, Daverio M, Palazzani L, Dimitriou D, Diez-Domingo J, Fons-Martinez J, et al. Digital tools in the informed consent process: a systematic review. *BMC Med Ethics*. 2021;22:1–10.

76. Markus AF, Kors JA, Rijnbeek PR. The role of explainability in creating trustworthy artificial intelligence for health care: A comprehensive survey of the terminology, design choices, and evaluation strategies. *J Biomed Inform*. 2021;113:103655.

77. Saeidnia HR, Hashemi Fotami SG, Lund B, Ghiasi N. Ethical Considerations in Artificial Intelligence Interventions for Mental Health and Well-Being: Ensuring Responsible Implementation and Impact. *Soc Sci*. 2024;13:381.

78. Smith A, Severn M. An Overview of Continuous Learning Artificial Intelligence-Enabled Medical Devices. *Canadian Journal of Health Technologies*. 2022;2.
79. Papagiannidis E, Mikalef P, Conboy K. Responsible artificial intelligence governance: A review and research framework. *The Journal of Strategic Information Systems*. 2025;34:101885.
80. Feng J, Phillips R V, Malenica I, Bishara A, Hubbard AE, Celi LA, et al. Clinical artificial intelligence quality improvement: towards continual monitoring and updating of AI algorithms in healthcare. *npj Digital Medicine* 2022 5:1. 2022;5:1–9.
81. Chustecki M. Benefits and Risks of AI in Health Care: Narrative Review. *Interact J Med Res* 2024;13:e53616 <https://www.i-jmr.org/2024/1/e53616>. 2024;13:e53616.
82. Ali G, Mijwil MM. Cybersecurity for Sustainable Smart Healthcare: State of the Art, Taxonomy, Mechanisms, and Essential Roles. *Mesopotamian Journal of CyberSecurity*. 2020. <https://doi.org/10.58496/MJCS/2024/006>.
83. Tsihrintzis GA, Virvou M, Washizaki H, Phillips-Wren G, Miltos M(, Alamaniotis), et al. Federated Learning: Navigating the Landscape of Collaborative Intelligence. *Electronics* 2024, Vol 13, Page 4744. 2024;13:4744.
84. Abbas SR, Abbas Z, Zahir A, Lee SW. Federated Learning in Smart Healthcare: A Comprehensive Review on Privacy, Security, and Predictive Analytics with IoT Integration. *Healthcare* 2024, Vol 12, Page 2587. 2024;12:2587.
85. Marey A, Arjmand P, Alerab ADS, Eslami MJ, Saad AM, Sanchez N, et al. Explainability, transparency and black box challenges of AI in radiology: impact on patient care in cardiovascular radiology. *Egyptian Journal of Radiology and Nuclear Medicine*. 2024;55:1–14.
86. Sadeghi Z, Alizadehsani R, CIFCI MA, Kausar S, Rehman R, Mahanta P, et al. A review of Explainable Artificial Intelligence in healthcare. *Computers and Electrical Engineering*. 2024;118.
87. Amann J, Blasimme A, Vayena E, Frey D, Madai VI. Explainability for artificial intelligence in healthcare: a multidisciplinary perspective. *BMC Med Inform Decis Mak*. 2020;20:1–9.
88. Gurevich E, El Hassan B, El Morr C. Equity within AI systems: What can health leaders expect? *Healthc Manage Forum*. 2022;36:119.
89. Hanna MG, Pantanowitz L, Jackson B, Palmer O, Visweswaran S, Pantanowitz J, et al. Ethical and Bias Considerations in Artificial Intelligence/Machine Learning. *Modern Pathology*. 2025;38:100686.
90. Ueda D, Kakinuma T, Fujita S, Kamagata K, Fushimi Y, Ito R, et al. Fairness of artificial intelligence in healthcare: review and recommendations. *Jpn J Radiol*. 2023;42:3.
91. Kiseleva A, Kotzinos D, De Hert P. Transparency of AI in Healthcare as a Multilayered System of Accountabilities: Between Legal Requirements and Technical Limitations. *Front Artif Intell*. 2022;5:879603.
92. Wilkinson MD, Dumontier M, Aalbersberg IJ, Appleton G, Axton M, Baak A, et al. The FAIR Guiding Principles for scientific data management and stewardship. *Scientific Data* 2016 3:1. 2016;3:1–9.
93. Gierend K, Waltemath D, Ganslandt T, Siegel F. Traceable Research Data Sharing in a German Medical Data Integration Center With FAIR (Findability, Accessibility, Interoperability, and Reusability)-Geared Provenance Implementation: Proof-of-Concept Study. *JMIR Form Res* 2023;7:e50027 <https://formative.jmir.org/2023/1/e50027>. 2023;7:e50027.
94. de Mello BH, Rigo SJ, da Costa CA, da Rosa Righi R, Donida B, Bez MR, et al. Semantic interoperability in health records standards: a systematic literature review.

Health Technol (Berl). 2022;12:255–72.

95. Hanna MG, Pantanowitz L, Jackson B, Palmer O, Visweswaran S, Pantanowitz J, et al. Ethical and Bias Considerations in Artificial Intelligence/Machine Learning. *Modern Pathology*. 2025;38:100686.

96. Willacott A. Developing an integrated national PROMs and PREMs platform for NHS Wales in *Advances in Patient Reported Outcomes: Integration and Innovation*. *Journal of Patient-Reported Outcomes* 2020 4:1. 2020;4 1: 27:1–16.

97. Kalluri M, Luppi F, Vancheri A, Vancheri C, Balestro E, Varone F, et al. Patient-reported outcomes and patient-reported outcome measures in interstitial lung disease: where to go from here ? *European Respiratory Review*. 2021;30 April.

98. Rudolph AE, Bazzi AR, Fish S. Ethical considerations and potential threats to validity for three methods commonly used to collect geographic information in studies among people who use drugs. *Addictive Behaviors*. 2016;84–90.

99. Kwame A, Petrucka PM. A literature-based study of patient-centered care and communication in nurse-patient interactions: barriers, facilitators, and the way forward. *BMC Nurs*. 2021;20:1–10.

100. Tilala MH, Chenchala PK, Choppadandi A, Kaur J, Naguri S, Saoji R, et al. Ethical Considerations in the Use of Artificial Intelligence and Machine Learning in Health Care: A Comprehensive Review. *Cureus*. 2024;16:e62443.

101. Ning Y, Teixayavong BSS S, Shang Y, Miao D, W Ting DS, Liu M, et al. Generative artificial intelligence and ethical considerations in health care: a scoping review and ethics checklist. *Lancet Digit Health*. 2024;6:e848–56.

102. Tang L, Li J, Fantus S. Medical artificial intelligence ethics: A systematic review of empirical studies. *Digit Health*. 2023;9.

103. Badawy W, Zinhom H, Shaban M. Navigating ethical considerations in the use of artificial intelligence for patient care: A systematic review. *Int Nurs Rev*. 2024. <https://doi.org/10.1111/INR.13059>.

104. Mennella C, Maniscalco U, De Pietro G, Esposito M. Ethical and regulatory challenges of AI technologies in healthcare: A narrative review. *Heliyon*. 2024;10:e26297.

105. Nasa P, Jain R, Juneja D. Delphi methodology in healthcare research: How to decide its appropriateness. *World J Methodol*. 2021;11:116.

106. Esmaeilzadeh P, Mirzaei T, Dharanikota S. Patients' Perceptions Toward Human–Artificial Intelligence Interaction in Health Care: Experimental Study. *J Med Internet Res*. 2021;23:e25856.

107. Bernasconi L, Şen S, Angerame L, Balyegisawa AP, Hong Yew Hui D, Hotter M, et al. Legal and ethical framework for global health information and biospecimen exchange- A n international perspective. *BMC Med Ethics*. 2020;21:1–8.

Supplementary Files