

Artificial Intelligence: Ethics in usability.**Dr Rishi Jerath**

Asstt Proff Commerce Government College Nalagarh.

Abstract *Artificial intelligence (AI) has profoundly changed and will continue to change our lives. AI is being applied in more and more fields and scenarios such as autonomous driving, medical care, media, finance, industrial robots, and internet services. The widespread application of AI and its deep integration with the economy and society have improved efficiency and produced benefits. At the same time, it will inevitably impact the existing social order and raise ethical concerns. Ethical issues, such as privacy leakage, discrimination, unemployment, and security risks, brought about by AI systems have caused great trouble to people. Therefore, AI ethics, which is a field related to the study of ethical issues in AI, has become not only an important research topic in academia, but also an important topic of common concern for individuals, organizations, countries, and society. This paper will give a comprehensive overview of this field by summarizing and analyzing the ethical risks and issues raised by AI, ethical guidelines and principles issued by different organizations, approaches for addressing ethical issues.*

Keywords: *Artificial Intelligence, Autonomous, Issues Industrial, Ethical, Principles.*

1. INTRODUCTION

There is nearly universal agreement among modern AI professionals that Artificial Intelligence falls short of human capabilities in some critical sense, even though AI algorithms have beaten humans in many specific domains such as chess. It has been suggested by some that as soon as AI researchers figure out how to do something, capability ceases to be regarded as intelligent—chess was considered the epitome of intelligence until Deep Blue won the world championship from Kasparov—but even these researchers agree that something important is missing from modern AI. A variety of examples of ethical challenges related to AI, organized into four key areas: design, process, use, and impact. “The design category includes decisions about what to build, how, and for whom,” they explain. “The process category includes decisions about how to support transparency and accountability through institutional design. The use category includes ways in which AI systems can be used and misused to cause harm to individuals or groups. Lastly, the impact category includes ways in which AI technologies result in broader social, political, psychological, and environmental impacts.”

Ethics Approach:

AI ethics are the set of guiding principles that stakeholders (from engineers to government officials) use to ensure artificial intelligence technology is developed and used responsibly. This means taking a safe, secure, humane, and environmentally friendly approach to AI. With its unique mandate, UNESCO has led the international effort to ensure that science and technology develop with strong ethical guardrails for decades.

Be it on genetic research, climate change, or scientific research, UNESCO has delivered global standards to maximize the benefits of the scientific discoveries, while minimizing the downside risks, ensuring they contribute to a more inclusive, sustainable, and peaceful world. It has also identified frontier challenges in areas such as the ethics of neuron technology, on climate engineering, and the internet of things.

AI and Work: Uses and Unique Features:

Current AIs constitute artificial narrow intelligence, or AIs that can undertake actions only within restricted domains, such as classifying pictures of cats . The “holy grail” of AI research is artificial general intelligence, or AIs that can perform at least as well as humans across the full range of intelligent activities. We focus only on narrow AI as it is already used across many diverse sectors, including in healthcare, judicial, educational, manufacturing, and military contexts, among many others (see. The established use of narrow AI also allows us to draw on practical examples to ground our assessment of its effects on meaningful work. While considering the possible implications of artificial general intelligence for meaningful work is important, and we discuss this in our future research directions, there remain persistent disagreements about when, if ever, it will be achieved. This makes it critical to examine the impacts of current AI capabilities on opportunities for meaningful work that are occurring now and in the near-term.

Past research demonstrates the dual positive and negative effects of technology upon aspects of meaningful work. For example, technology use can up skill workers and enhance their autonomy, but it can also deskill and serve to control the, with meaningfulness generally elevated in the former case. Technology’s positive effects can also help individuals confirm pre-existing notions of meaningful work, but its negative outcomes can require them to re-interpret and adjust those meanings as the technology’s dual effects are realized , for example by providing on-demand connection to work but heightening distraction from other responsibilities. Such dual effects remain evident in advancing forms of technology, such as workplace robotics that offer both benefits and threats to meaningful human work .

The Ethics of Meaningful Work and Ethical AI

Recent philosophical discussions of meaningfulness tend to focus on what makes life itself, or the activities and relationships that compose a well-lived life, meaningful. The paradigm of meaningless work is Sisyphus, who is condemned as punishment to repeatedly roll a rock to the top of a mount. Sisyphus’ work is boring, repetitive, simple, does not benefit others, and is not freely .By implication, meaningful work should be engaging, varied, require the use of complex skills, benefit others, and be freely chosen. This emphasizes two aspects of meaningfulness that Wolf calls subjective (do you experience work as meaningful?) and objective (is the work actually meaningful?) elements. As we take meaningful work to be “personally significant and worthwhile” our definition is inclusive of these subjective (it is personally significant) and objective (it is worthwhile) aspects. The Ethical Implications of Meaningful Work: Why is it Ethically Important?

AI and research ethics

Researchers are faced with several ethical quandaries while navigating research projects. They are urged to safeguard human research participants' protection when working with human research participants. However, it is not always simple to balance the common good (i.e., develop solutions for the wider population) and the individual interest (i.e., research participants' safety). Researchers are responsible for anticipating and preventing risks from harming participants while advancing scientific knowledge, which requires maintaining an adequate risk-benefit ratio. With AI's fast growth, another set of issues is added to the existing ones: data governance, consent, responsibility, justice, transparency, privacy, safety, reliability, and more,. This section will describe the views on current guidelines to regulate AI, key principles and ethical approaches, and the main issues. In the current climate, we expect

continuity on the following concepts: responsibility, explain ability, validity, transparency, informed.

. Moral status and rights

While guidelines and norms are shifting to fit AI standards, many questions on moral status and rights are raised to adapt to this new reality. Authors argue that we cannot assign moral agency to AI and robots. There are multiple reasons for it: robots do not seem capable of solving problems ethically, AI's lack of explanation regarding its generated results, and the absence of willingness to choose which might impact decision making in research ethics.

Rights are attributed to different living entities. For instance, in the EU, the law protects animals as sentient living organisms and unique tangible goods. Their legal status also obliges researchers not to harm animals during research projects, making us question the status and rights we should assign AIS or robots. Indeed, Miller pointed out that having a machine at one's disposal raises questions on human-machine relationships and the hierarchical power it might induce.

Key principles and norms of AI systems in research ethics

We have seen that the lexicon and the language used invoke both classical theories and contextualization of AI ethics benchmarks within the practices and ethos of research ethics.

Ethical approaches in terms of AI research ethics

The literature invoked the following classic theories: the Kantian-inspired model, utilitarianism, principlism (autonomy, beneficence, justice, and non-maleficence), and the precautionary principle.

Responsibility in AI research ethics

Public education and ethical training implementation could help governments spread awareness and sensitize people regarding research ethics in AI. Accountability of AI regulation and decision-making should not strictly fall into stakeholders' hands but also be based on solid legal grounds. Digital mental health apps and other institutions will now be attributed responsibilities that have usually been acclaimed to professionals or researchers using the technology (i.e., decision-making, providing users with enough tools to understand and use products, being able to help when needed, etc. Scientists and AI developers must not throw caution to the wind regarding the possibility that biased algorithms could be fed to AI models. Clinicians will have to tactfully manage to inform patients of the results generated by machine learning (ML) models while considering their risk of error and bias. It is still vague to attribute responsibility to specific actors. However, it is necessary to have different groups work together to tackle the problem. Some consider that validity, explainability, and open-source AI systems are some of the defining points that lead to responsibility. With the advancement of technologies and its gain of interest, the sense of social responsibility also increased. Indeed, every actor must contribute to making sure that these novel technologies are developed and used in an ethical matter.

Explainability and validity

An important issue with AIS usually raised is the explainability of results. Deep learning (DL) is another type of ML with more extensive algorithms that encloses data with a broader array

of interpretations. This makes it harder to explain how DL and AI models reached a particular conclusion. This poses transparency issues that are challenging to participants.

Since AI is known for its 'black-box' aspect, where results are difficult to justify, it is difficult to fully validate a model with certainty. Deciding to monitor research participants closely could help validate results which, in theory, would bring more accurate results. However, close monitoring could also have a negative effect by influencing participants' decisions based on whether they mind being monitored or not. This event could, as a result, produce more inaccurate results. Furthermore, it could be more challenging in certain contexts to promote validity when journals and funding bodies favor new and innovative studies over ethical research on AI, even if the latter is being promoted.

Transparency and informed consent

According to the White House Office of Science and Technology Policy (OSTP), transparency would help solve many ethical issues. Transparency allows research participants to be aware of a study's different outlooks and comprehend them. The same goes for new device users. AI models (i.e., products, services, apps, sensor-equipped wearable systems, etc.) produce a great deal of data that does not always come from consenting users. Furthermore, AI's black-box imposes a challenge to obtain informed consent since the lack of explainability of AI-generated results might not allow participants to have enough information to give out their informed consent. Thus, it is essential to make consent forms easy to understand for the targeted audience.

However, the requirement to get informed consent could lead to other less desirable implications. Some argue that requiring authorization for all data, especially studies that hold a vast set of data, might lead to data bias and a decrease in data quality because it only entices a specific group of people to give out consent which leaves out a significant part of the population.

Privacy

While the levels of privacy differ from one scholar to another, the concept of privacy remains a fundamental value to human beings. Through AI and robotics, data can be seen as attractive commodities which could compromise privacy. Researchers are responsible for keeping participants unidentifiable while using their data. However, data collected from many sources can induce a higher risk of identifying people. While pursuing their research study, ML researchers still struggle to comply with privacy guidelines.

Data protection

According to a study, most people do not think data protection is an issue. One reason to explain this phenomenon is that people might not fully grasp the magnitude of its impact. Indeed, the effect could be very harmful to some people. For instance, data found about a person could decrease their chances of employment or even of getting insurance. Instead of focusing on data minimization, data protection should be prioritized to ensure ML models get the most relevant data, ensuring data quality while maintaining privacy. Another point worth mentioning is that the GDPR allows the reuse of personal data for research purposes, which might allow companies who wish to pursue commercial research to bypass certain ethical requirements.

Privacy vs. science advancement dilemmas

Some technology-based studies face a dichotomy between safeguarding participants' data and making scientific advancements. This does not always come easily since ensuring privacy can

compromise data quality, while studies with more accurate data usually lead to riskier privacy settings. Indeed, with new data collection methods in public and digital environments, consent and transparency might be overlooked for better research results.

Key issues of the current state of AI-specific RE guidelines

Many difficulties have arisen with the soaring evolution of AI. There has been a gap between research ethics and AI research, inconsistent standards regarding AI regulation and guidelines, and a lack of knowledge and training in these new technologies has been widely noticed. Medical researchers are more familiar with research ethics than computer science researchers and technologists. This shows a disparity in knowledge between different fields.

With new technologies comes the difficulty in assessing them. Research helps follow AI's progress and ensures it does so responsibly and ethically. Unfortunately, applied and research ethics are not always in sync. AI standards mostly rely on ethical values rather than concrete normative and legal regulations, which have become insufficient. The societal aspects of AI are more discussed amongst researchers than the ethics part of research.

Many countries have taken the initiative to regulate AI using ethical standards. However, guidelines vary from one region to another. It has become a strenuous task to establish a consensus of strategies, turn principles into laws, and make them practical. It does not only come down to countries that have differing points of views but journals as well. Indeed, validation for an AI research project publication could differ from one journal to another. Even though ethical, legal, and social implications (ELSI) are used to help oversee AI, regulations and AI-specific guidelines remain

When research ethics guidelines are applied to AI

While there is a usual emphasis that is being made on ethical approbation for research projects, there are other projects that are not required to follow an ethics guideline. In the United Kingdom, some research projects do not require ethics approval (i.e., social media data, geolocation data, anonymous secondary health data with an agreement). A study highlighted that most papers gathered that used available data from social media did not have an ethical approbation. Some technology-based research projects ask for consent from their participants but skip requesting ethical approval from a committee. Some non-clinical research projects are exempt from an ethics evaluation. Tools do not always undergo robust testing before validation either. Of course, ethical evaluation remains essential in multiple other settings: when minors or people lacking capacity to make an informed decision are involved, when users are recognizable, when researchers seek users' data directly, when clinical data or applications are used, etc.

Research ethics board

Historically REBs have focused on protecting human participants in research (e.g., therapeutic, nursing, psychological, or social research) from complying with the requirements of funding or federal agencies like NIH or FDA. This approach has continued, and in many countries, REBs are fundamentally essential to ensure that research involving human participants is conducted in compliance with ethics guidelines and national and international regulation

Roles of REB

The primary goal of REBs focuses on reviewing and overseeing research to provide the necessary protection for research participants. REBs consist of groups of experts and stakeholders (clinicians, scientists, community members) who review research protocols with

an eye toward ethical concerns. They ensure that protocols comply with regulatory guidelines and can withhold approval until such matters have been addressed. Also, they were designed to play an anticipatory role, predicting what risks might arise within research and ironing out ethical issues before they appeared. Accordingly, REBs aim to assess whether the proposed research project meets specific ethical standards regarding the foreseeable impacts on human subjects. However, REBs are less concerned with the broader consequences of research and its downstream applications. Instead, they focus on the direct effects on human subjects during or after the research process. Within their established jurisdiction, REBs can develop a review process independently. Considering the specific context of AI research, REBs would aim to mitigate the risks of potential harm possibly caused by technology. This could be done by reviewing scientific questions relating to the origin and quality of the data, algorithms, and artificial intelligence; confirming the validation steps conducted to ensure the prediction models work; requesting further validation to be carried out if required.

. Scope and approaches

AI technologies are rapidly changing health research; these mutations might lead to significant gaps in REB oversight. Some authors who analyzed these challenges suggest an adaptive scope and approach. To achieve an AI-appropriate research ethics review, it is necessary to clearly define the thresholds and characteristics of cardinal research ethics considerations, including what constitutes a “*human participant*”, what is a *treatment*, what is a *benefit*, what is a *risk*, what is considered a publicly available information, what is considered an intervention in the public domain, what is a medical data, but also what is AI research”.

There is an urgent need to tailor the technology and its development, evaluation, and use contexts (i.e., digital mental health). Health research involving AI features requires intersectoral and interdisciplinary participatory efforts to develop dynamic, adaptive, and relevant normative guidance. It also requires practice navigating the ethical, legal, and social complexities of patient data collection, sharing, analysis, interpretation, and transfer for decision-making in a natural context. Also, these studies imply multi-stakeholder participation (such as regulatory actors, education, and social media).

This diversity of actors seems to be a key aspect in this case. Still, it requires transparent, inclusive, and transferable normative guidance and norms to ensure that all understand each other and meet the normative demands regarding research ethics. Furthermore, bringing together diverse stakeholders and experts is worthwhile, especially when the impact of research can be significant, difficult to foresee, and unlikely to be understood by any single expert, as with AI-driven medical research. In this stake, several factors are beneficial to promote cooperation between academic research and industry: inter-organizational trust, collaboration experience, and the breadth of interaction channels. Partnership strategies like collaborative research, knowledge transfer, and research support may be essential to embolden this in much broader terms than strict technology transfer.

. AI research ethics, practices, and governance oversight

According to the results of our review, REBs must assess the following seven considerations of importance during AI research ethics review: (1) Informed consent, (2) benefit-risks assessment, (3) safety and security, (4) validity and effectiveness, (5) user-centric approach and design, (6) transparency. In the literature, some authors have pointed out specific questions about considerations REBs should be aware of.

Informed consent

Some authors argue that the priority might be to consider whether predictions from a specific machine learning model are appropriate for informing decisions about a particular intervention. Others advocate carefully constructing the planned interventions so research participants can understand them.

The extent to which researchers should provide extensive information to participants is not evident among stakeholders. So far, research suggests that there is no clear consensus among patients on whether they would want to know this kind of information about themselves. Hence the question remains whether patients want to see if they are at risk, mainly if they cannot be told why, as factors included in machine learning models generally cannot be interpreted as having a causal impact on outcomes. Therefore, sharing information from an uninterruptable model may adversely affect a patient's perception of their illness, confuse them, and immediate concerns about transparency.

Benefits/risks assessment

The analysis of harms and potential benefits is critical when assessing human research. REBs are well concerned with this assessment to prevent unnecessary risks and promote benefits. Considerations of the potential benefits and harms to patient-participants are necessary for future clinical research, and REBs are optimally positioned to perform this assessment. Additional considerations like benefit/risk ratio or effectiveness and the systematic process described previously are necessary. Risk assessments could have a considerable impact in research involving mobile devices or robotics because preventive action and safety measures may be required in the case of imminent risks. Thus, REB risk assessment seems very important.

Approaching AI research ethics through user-centered design can represent an interesting avenue to understand better how REB can conduct risk/benefices assessment. For researchers, involving users in the design of AI research is likely to promote better research outcomes. Hence, this can be reached by investigating how AI research is actually meeting users' needs and how this may generate intended and unintended impacts on them. Indeed there is insufficient reason to believe that AI research will produce positive benefits unless it is evaluated with a focus on patients and situated in the context of clinical decision-making. Consequently, REBs might focus on the broader societal impact of this research.

Safety and security

Safety and security are significant concerns for AI and robotics, and their assessment may rely on end-users' perspectives. To address the safety issue, it is not sufficient for robotics researchers to say that their robot is safe based on literature and experimental tests. It is crucial to find out about the perception and opinions of end-users of robots and other stakeholders. Testing technology in real-life scenarios is vital for identifying and adequately assessing technology's risks, anticipating unforeseen problems, and clarifying effective monitoring mechanisms. On the other hand, there is a potential risk that an AIS misleads the user in realizing a legal act.

Validity and effectiveness

Validity is a crucial consideration and one on which there is consensus to appreciate the normative implications of AI technologies. To this end, it is necessary for research ethics that researchers' protocols be explicit about many elements and describe their validation model and

performance metrics in a way that allows for assessment of the clinical applicability of their developing technology. In addition, in terms of validity, simulative models have yet to be appropriately compared with standard medical research models (including *in vitro*, *in vivo*, and clinical models) to ensure they are correctly validated and effective. Considering many red flags raised in recent years, AI systems may not work equally well with all sub-populations (racial, ethnic, etc.). Therefore, AI systems must be validated for different subpopulations of patients.

Demonstration of value is essential to ensure the scientific validity of the claims made for technology but also to attest to the proven effectiveness once deployed in a real-world setting and the social utility of a technology. When conducting a trial for the given AI system, the main interest should be to assess its overall reliability, while the interaction with the clinician might be less critical

Transparency

Transparency entails understanding how technology behaves and establishing thresholds for permissible (and impermissible) usages of AI-driven health research. Transparency requires clarifying the reasons and rationales for the technology's design, operation, and impacts. Identified risks should be accompanied by detailed measures intended to avoid, reduce, or eliminate the risks. The efficiency of such efforts should be assessed upstream and downstream as part of the quality management process. As far as possible, testing methods, data, and assessment results should be public. Transparent communication is essential to make research participants, as well as future users aware of the technology's logic and functioning.

The analysis also showed that data bias is a flagrant problem whether AI is used or not and that this discriminatory component should be taken care of to avoid emphasizing the problem with AI. Informed consent is another value that REBs prioritize and will have to be adapted to AI because new information might have to be disclosed to participants. Safety and security are always essential to consider. However, other measures will be implemented with AI to ensure that participants are not set in danger. One of the main aspects of AI is data sharing and the risk that this might breach participants' privacy. The methods put in place now might not be suitable for AI's fast evolution. The questions of justice, equality, and fairness that have not been resolved in our current society will also have to be instigated in the AI era. Finally, the importance of validity was raised numerous times. Unfortunately, REBs do not have the right tools to evaluate AI. It will be necessary for AI to meet the population's needs. Furthermore, definitions of specific values and principles that REBs usually respond to will have to be reviewed and adapted according to AI.

2. LIMITATIONS AND CHALLENGES

Our results point to several discrepancies between the critical considerations for AI research ethics and REB review of health research and AI/ML data.

Consent forms

According to our review, there is a disproportionate focus on consent before other ethical issues. Authors argue that the big piece the REBs ask for relies on consent, not the AI aspect of the project. This finding suggests that narrowing AI research ethics around consent concerns remains problematic. In some stances, the disproportionate focus on consent, along with the importance REBs place on consent forms and participant information sheets, has settled how research ethics is defined, e.g., viewed as a proxy for ethics best practice, or in some cases, as an ethics panacea.

Safety, security, and validity

Authors report a lack of knowledge for safety review. It appears clear that REBs may not have the experience or expertise to conduct a risk assessment to evaluate the probability or magnitude of potential harm. Similarly, the training data used to inform the algorithm development are often not considered to qualify as human subjects research, which – even in a regulated environment – makes a prospective review for safety potentially unavailable.

On the other hand, REBs lack appropriate assessing processes for assessing whether AI systems are valid, effective, and apposite. The requirement to evaluate the evidence of effectiveness adds to a range of other considerations with which REBs must deal (i.e., the protection of participants and the fairness in the distribution of benefits and burdens). Therefore, there is still much to be done to equip REBs to evaluate the effectiveness of AI technologies, interventions, and research.

Privacy and confidentiality

Researchers point to a disproportionate focus on data privacy and governance before other ethical issues in medical health research with AI tools. Focus on privacy and data governance issues warrants further attention, as privacy issues may overshadow other issues. Indeed it seems problematic and led to a narrowing of ethics and responsibility debates being perpetuated throughout the ethics ecosystem, often at the expense of other ethical issues, such as questions around justice and fairness. REBs appear to be less concerned about the results themselves. One explained that when reviewing their AI-associated research ethics applications, REBs focus more on questions of data privacy than other ethical issues, such as those related to the research and the research finding. Others painted a similar picture of how data governance issues were a centralized focus when discussing their interactions with their REB. According to these stakeholders, REBs focus less on the actual algorithm than how the data is handled, and the issue remains around data access and not about the software.

Governance, oversight, and process

Lack of expertise appears to be a significant concern in our results. Indeed, even when there is oversight from a research ethics committee, authors observe that REB members often lack the experience or confidence regarding particular issues associated with digital research.

Some authors advocate that ML researchers should complement the membership of REBs since they are better situated to evaluate the specific risks and potential unintended harms linked to the methodology of ML. On the other hand, REBs should be empowered to fulfill their role in protecting the interests of patients and participants and enable the ethical translation of healthcare ML. However, we can notice that researchers expressed different views about REBs' expertise. While most acknowledged a lack of AI-specific proficiency, for many, this remains straightforward because the ethical issues of their AI research were no exceptional compared to other ethics issues raised by “big data”.

Limits of process and regulation are another concern faced by REBs, including a lack of consistency in decision-making within and across REBs, a lack of transparency, poor representation of the participants and public they are meant to represent, insufficient training, and a lack of measures to examine their effectiveness. There are several opinions on the need for and the effectiveness of REBs, with critics lamenting excessive bureaucracy, lack of reliability, inefficiency, and, importantly, high variance in outcomes. To address the existing gap of knowledge between different fields, training could be used to help rebalance this and ensure sufficient expertise for all research experts to pursue responsible innovation

Researchers described the lack of standards and regulations for governing AI at the level of societal impact; the way that ethics committees in institutions work is still acceptable. But there is a need for another level of thinking that combines everything and does not look at one project simultaneously.

Finally, researchers have acknowledged the lack of ethical guidance, and some REBs report feeling ill-equipped to keep pace with rapidly changing technologies used in research.

Stakeholder perceptions and engagement

Researchers' perspectives on AI research ethics may vary. While some claim that researchers often take action to counteract the adverse outcomes created by their research projects, others promulgate that researchers do not always notice these outcomes. When the latter occurs, researchers are pressed to find solutions to deal with those outcomes.

Furthermore, researchers are expected to engage more in AI research ethics. Researchers must demonstrate cooperation with certain institutions (i.e., industries and governments). Researchers are responsible for ensuring that their research project is conducted responsibly by considering participants' needs. Usually, research ethics consist of different researchers coming from multidisciplinary fields who are better equipped to answer further ethical and societal questions. However, there could be a clash of interests between parties while setting goals for a research project.

A lot of the time, different stakeholders do not necessarily understand other groups' realities. Therefore, research is vital to ensure that stakeholders can understand one another and be in the same scheme of things. This will help advance AI research ethics.

Responsibility for ensuring a responsible utilization of AI lies within various groups of stakeholders.

Many others, such as the private sector, can be added to the list. Studies have shown that private companies' main interest is profit over improving health with the data collected using AI. Another problematic element with the private sector: they do not often fall under the regulation of ethical oversight boards, which means that AI systems or robots that come from private companies do not necessarily follow an accepted ethical guideline. This goes beyond ethical research concerns.

Key practices and processes for AI research

REBs may face new challenges in the context of research involving AI tools. Authors are calling for specific oversight mechanisms, especially for medical research projects.

. Avoid bias in AI data

While AI tools provide new opportunities to enhance medical health research, there is an emerging consensus among stakeholders regarding bias concerns in AI data, particularly in clinical trials. Since bias can worsen pre-existing disparities, researchers should proactively target a wide range of participants to establish sufficient evidence of an AI system's clinical benefit across different populations. To mitigate selection bias, REBs may require randomization in AI clinical trials. To achieve this, researchers must start by collecting more and better data from social minority groups. Also, awareness of biases concerns should be taken into account in the validation phase, where the performance of the AI system gets measured in a benchmark data set. Hence it is crucial to test AI systems for different subpopulations. Therefore, affirmative action in recruiting research participants in AI RCTs deems us ethically

permissible. However, authors reported that stakeholders might encounter challenges accessing needed data in a context where severe legal constraints are imposed on sharing medical data.

Attention to vulnerable populations

Vulnerable populations require excellent protection against risks they may face in research.

When involving vulnerable populations, such as those with a mental health diagnosis, in AI medical health research, additional precautions should be considered to ensure that those involved in the study are duly protected from harm – including stigma and economic and legal implications. In addition, it is essential to consider whether access barriers might exclude some people.

. Diversity, inclusion, and fairness

Another issue, which needs to be raised when considering critical practices and scope in AI research, relates to fair representation, diversity, and inclusion. According to Grote, one should explore concerns for the distribution of research participants and representatives for the state, country, or even world region in which the AI system gets tested. Here the author advocates if we should instead aim for a parity distribution of different gender, racial and ethnic groups. Hence, he raised several questions to support the reflection of REB on diversity, inclusion, and fairness issues: How should the reference classes for the different subpopulations be determined? Also, what conditions must be met for fair subject selection in AI RCTs? And finally, when, if ever, is it morally justifiable to randomize research participants in medical AI trials?

Guidance to assess ethical issues in research involving robotics

The aging population and scarcity of health resources are significant challenges healthcare systems face today. Consequently, people with disabilities, especially elders with cognitive and mental impairments, are the most affected. The evolving field of research with assistive robots may be useful in providing care and assistance to these people. However, robotics research requires specific guidance when participants have physical or cognitive impairments. Indeed particular challenges are related to informed consent, confidentiality, and participant rights. According to some authors, REBs should ask several questions to address these issues: Is the research project expected to enhance the quality of care for the research participants? What is/are the ethical issue/s illustrated in this study? What are the facts? Is any important information not available in the research? Who are the stakeholders? Which course of action best fits with the recommendations and requirements set out in the “Ethical Considerations” section of the study protocol? How can that course of action be implemented in practice? Could the ethical issue/s presented in the case be prevented? If so, how?

Which ethical and social issues may neuron robotics raise, and are mechanisms currently implemented sufficiently to identify and address these? Is the notion that we may analyze, understand and reproduce what makes us human rooted in something other than

3. UNDERSTANDING OF THE PROCESS BEHIND AI/ML DATA

A good understanding of the process behind AI/ML tools might be of interest to REBs when assessing the risk/benefit ratio of medical research involving AI. However, there seems to be a lack of awareness of how AI researchers gain results. Authors argue that it would not be impossible to induce perception about the external environment in the neuron culture and to interpret the signals from the neuron culture as motor commands without a basic understanding of this neural code. Indeed, when using digital health technologies, the first step is to ask

whether the tools, be they apps or sensors, or AI applied to large data sets, have demonstrated value for outcomes. One should ask whether they are clinically effective, or if they measure what they purport to measure (validity) consistently (reliability), and finally, if these innovations also improve access to those at the highest risk of health disparities.

Indeed, the ethical issues of AI research raise major questions within the literature. What may seem surprising at first sight is that the body of literature is still relatively small and appears to be in an embryonic state regarding the ethics of the development and use of AI (outside the scope of academic research). The literature is thus more concerned with the broad questions of what constitutes research ethics in AI-specific research and with pointing out the gaps in normative guidelines, procedures, and infrastructures adapted to the oversight of responsible and ethical research in AI. Perhaps unsurprisingly, most of the questions related to study within the health sector. This is to be expected given the ascendancy of health in general within the research ethics field. Thus, most considerations relate to applied health research, the implications for human participants (whether in digital health issues, research protocols, or interactions with different forms of robots), and whether projects should be subject to ethics review.

Specifically in AI-specific research ethics, interestingly, traditional issues of participant protection (including confidentiality, consent, and autonomy in general) and research involving digital technologies intersect and are furthered by the uses of AI. Indeed, as AI requires big data and behaves very distinctly from other technologies, the primary considerations raised by the body of literature studied were predominantly classical in AI ethics but contextualized and exacerbated within research ethics practices. For instance, one of the most prevalent ethical considerations raised and discussed was privacy and the new challenges regarding the massive amount of data collected and its use. If a breach of confidentiality were to happen or data collection would lead to discovering further information, this would raise the possibility of harming individuals. In addition, informed consent was widely mentioned and focused on transparency and explainability when the issues were AI-specific. Indeed, AI's black-box issue of explainability was raised many times. This is a challenge because it is not always easy to justify the results generated by AI. This then poses a problem with transparency. Indeed, participants expect to have the necessary information relevant to the trial to make an informed and conscious decision regarding their participation. Not having adequate knowledge to share with participants might not align with informed consent.

Furthermore, another principle was brought up many times, which was responsibility. Responsibility is shared chiefly between the researcher and the participant. Now that AI is added to the equation, it has become harder to determine who strictly should be held accountable for the occurrence of certain events (i.e., data error) and in what context. While shared responsibility is an idea many share and wish to implement, it is not easy. Indeed, as seen in many stakeholders (e.g., lawmakers, AI developers, AI users) may participate in responsibility sharing. However, much work will have to be put into finding a fair way to share responsibility between each party involved.

Our results have implications mainly on three levels. Indeed, AI-specific implications for research ethics is first addressed followed by REBs who take on these challenges.

AI-specific implications for research ethics:

The issues raised by AI are eminently global. It is interesting to see in the articles presented in the scoping review that researchers in different countries are asking questions colored by the jurisdictional, social, and normative context in which the authors work. However, there appears

to be heterogeneity in the advancement of AI research ethics thinking; this is particularly evident in the progress of research ethics initiatives within countries. A striking finding is that very little development has been done regarding AI-specific standards and guidelines to frame and support research ethics worldwide.

At this point, the literature does not discuss the content of norms and their application to AI research. Instead, it makes initial observations about AI's issues and challenges to research ethics. In this sense, it is possible to see that the authors indicate new challenges posed by the emergence of AI in research ethics. AI makes many principles more challenging to assess (it seems quite difficult to use the current guidelines to balance the risks and benefits). One example is that it has become unclear which level of transparency is adequate. AI validity, on the other hand, is not always done in an optimal manner throughout AI's lifecycle. Accountability remains a continuing issue since it is still unclear who to hold accountable and to what extent with AI in play. In addition, AI is also known to amplify certain traditional issues in research ethics. For example, AI blurs the notion of free and informed consent since the information a patient or participant needs regarding AI is yet to be determined. Privacy's getting harder to manage because it has become possible with AI to identify individuals by analyzing all the data available, even after the identification. Data bias is another leading example where AI would not necessarily detect data bias it's being fed but could also generate more biased results.

Interestingly, the very distinction between the new AI-related issues and the old, amplified ones is still not entirely clear to researchers. For instance, while AI is quickly targeted for generating biased results, the source of the problem could come from biased data fed to AI. Another issue is the lack of robustness, where it is challenging to rely entirely on AI to always give accurate results. However, this issue is also found in human-based decision-making. Thus, the most efficient use of AI could depend on context. The final decision could be reserved for humans limiting AI's role as an assistive tool. Therefore, drawing a picture of what is new and less so is difficult. However, there is no doubt that AI is disrupting the field of research ethics, its processes, practices, and standards. This also points to the fact that no AI-specific Research Ethics Guidelines can help give a sense of how best to evaluate AI in a compatible way with RE guidance.

Another observation is that research ethics (and a fortiori research ethics committees) are very limited in their approach to AI development and research. This means that research ethics only comes into play at a specific point in developing AI technologies, interventions, and knowledge, i.e., after developing an AIS and before its implementation in a real context. Thus, research ethics, understood as it has been developed in most countries, focuses on what happens within public organizations and when human participants are involved. This excludes technological developments developed by industry and does not require ethical certification. Therefore, the vast majority of AIS outside the health and social services sector will not be subject to research ethics reviews, such as data found in social media or geolocation. But even within the health sector, AIS that do not directly interact with patients could largely be excluded from the scope of research ethics and the mandate of REBs. This makes the field of AI research ethics very new and small compared to responsible AI innovation.

Means for REBs

No author seems to indicate that REBs are prepared and equipped to evaluate research projects involving AI as rigorously, confidently, and consistently as for more traditional research protocols (i.e., not involving AI). One paper from expressively indicates that the current REB

model should be replicated in the private sector to help oversee and guide AI development. Arguably the call is mostly about adding an appraising actor to private sector technology developments than praising REBs for their mastery and competence in AI research ethics review. Yet, it still holds a relatively positive perception of the current readiness and relevance of REBs to research ethics. This may also reflect a lack of awareness (from uninformed stakeholders) of the limitations faced by REBs, which on paper can probably be seen as being able to evaluate research protocols involving AI and other projects. This is, however, disputed or refuted by the rest of the literature studied.

The bulk of the body of literature reviewed was more circumspect about the capacity of REBs. Not that they are not competent, but rather that they do not have the tools to start with a normative framework relevant to AI research, conceptually rigorous and comprehensive, and performative and appropriate to the mandates, processes, and practices of REBs. Over the last several decades, REBs have primarily relied on somewhat comprehensive and, to some extent, harmonized, regulations and sets of frameworks to inform and guide their ethical evaluation. The lack therefore, REBs face new challenges without any tools to support them with their decisions on AI dilemmas. The authors of our body of literature thus seem to indicate a higher expectation on all stakeholders to find solutions to address the specificities and challenges of AI in research ethics.

One of the first points is quite simple: determining when research involving AI should be subject to research ethics review. This simple point and observation is not consensual. Then, we can raise some serious concerns about the current mandate of REBs and the ability to evaluate AI with their current means and framework. Not only are they missing clear guidelines to do any kind of standard assessments on AI in research ethics, but they are also missing clearly defined roles on their account. Indeed, should their role be extended to look not just at research but also at the use of downstream technology? Or does this require another ethics oversight body that would look more at the technology in real life? This raises the question of how a lifecycle evaluative process can best be structured and how a continuum of evaluation can be developed that is adapted to this adaptive technology.

New research avenues

After looking at the heterogeneity of norms and regulations regarding AI in different countries, there should be an interest in initiating an international comparative analysis. The aim would be to investigate how REBs have adapted their practices in evaluating AI-based research projects without much input and support from norms. This analysis could raise many questions (i.e., could there be key issues that are impossible to universalize?).

The scope and approach of ethics review by REBs must be revisited in light of the specificities of research using AI:

The primary considerations we discuss above raise new challenges on the scope and approaches of REB practices when reviewing research with AI. Furthermore, applications developed within the research framework often rely on a population-based system, leading REBs to question whether their assessment should focus on a systematic individual approach rather than societal considerations and their underlying foundations.

However, AI research is still emerging, underlining the difficulties of completing such a debate. Finally, one can wonder about the importance of current guidelines in AI within the process of ethical evaluation by the REBs. Should this reflection be limited only to the REBs? Or should it include other actors meaning scientists or civil society?

AI ethics is not limited to research. While it is less discussed, AI ethics raises many existential questions. Dynamics such as the patient-physician relationship will have to adapt to a new reality (With human tasks being delegated to AI, notions of personhood autonomy (), and human status in our society (are threatened. This leads to delving into the question of “what it is to be human?”. Robots used in therapies aimed to care for patients (i.e., autistic children) could induce attachment issues and other psychological impacts . This projects another issue: AI overreliance, a similar problem brought up by current technological tools (i.e., cell phones)

. Updating and adapting processes in ethics committees

AI ethics is still an emerging field. The REBs ensure the application of ethical frameworks, laws, and regulations. Our results suggest that research in AI involves complex issues that emerge around these new research strategies and methodologies using technologies such as computer science, mathematics, or digital technology. Thus, REBs' concerns remain on recognizing and assessing ethical issues that arise from these studies and adapting to rapid changes in this emerging field.

In research ethics, respect for a person's dignity is essential. In several normative frameworks, i.e., the TCPS in Canada, it means respect for persons, concern for wellbeing, and justice. In AI research, REBs might need to reassess the notion of consent or the participant's place in the study. As with all research, REBs must ensure informed consent. However, there does not seem to be a clear consensus on the standard for providing informed consent in AI research. For example, REBs should consider the issue of AI's interpretability in a research consent form; to translate transparent and intelligible results.

Another issue that REBs consider here is the role of participants in AI research. Indeed, active participant involvement is not always necessary in AI research to complete the data collection to meet the research objectives. It is often the case when data collection is completed from connected digital devices or by querying databases. However, the consequences amplified the phenomenon of dematerialization of research participation while facilitating data circulation.

Furthermore, AI research and the use of these new technologies call on REBs to be aware of the changes this implies for the research participant, particularly concerns such as the continuous consent process, management of withdrawal, or the duration of participation in the research.

While protecting the individual participant takes center stage in the evaluation of REBs, research with AI may focus more on using data obtained from databases held by governments, private bodies, institutions, or academics. In this context, should concerns for societal wellbeing prevail over the wellbeing of the individual? There does not appear to be a clear consensus on what principles should be invoked to address this concern..

4. LIMITATIONS

The focal point of AI evaluation was often about privacy protection and data governance, not AI's ethics. While data protection and governance are massively essential issues, it should be equally important to investigate AI issues, not to leave out concerns that should be dealt with, such as AI validity, explainability, and transparency. In addition, FAIR and the ethics of care, which are starting to become standard approaches in the field, were not invoked in the articles to inform AI ethics in research. This might be due to the study's lack of literature on AI ethics compared to research ethics in general.

Another limitation worth outlining is that our final sample mainly reflected the reality and issues found in healthcare, despite having a scoping review open to all fields using AI. This could be due to the fact that AI is becoming more prominent in the field of healthcare (The field is also often linked to the development and presence of research ethics boards (Having healthcare outshine the rest of the fields in our sample could also be attributed to research ethics mostly stemming from multiple medical research incidents throughout history).

Furthermore, throughout the studied articles, few to none mentioned countries were non-affluent. This poses concerns about widening disparities between developed and developing countries. Therefore, it is vital to acknowledge the asymmetry of legislative and societal norms between countries to better serve their needs and avoid colonized practices.

Finally, this topic lacks maturity. This study primarily shows that REBs cannot find guidance from the literature. Indeed, there is a scarcity of findings in the literature regarding recommendations and practices to adopt in research using AI. There are even fewer findings that specifically aim to equip REBs. Reported suggestions are often about privileged behavior that governments or researchers should adopt rather than establishing the proper criteria REBs should follow during their assessments. Therefore, this study does not lead to findings directly applicable to REBs practice and should not be used as a tool for REBs.

5. CONCLUSION

Every field has its ethical challenges and needs. The results in this article have shown this reality. Indeed, we've navigated through some of AI ethics general issues before investigating AI ethics research-specific issues. This led us to discern what research ethics boards focus on more adeptly during their evaluations and the limits imposed on them when evaluating AI ethics in research. While AI is a promising field to explore and invest in, many caveats force us to develop a better understanding of these systems. With AI's development, many societal challenges will come our way, whether they are current ongoing issues, new AI-specific ones, or those that remain unknown to us. Ethical reflections are taking a step forward while normative guidelines adaptation to AI's reality is still dawdling. This impacts REBs and most stakeholders involved with AI. However, throughout the literature, many suggestions and recommendations were provided. This could allow us to build a framework with a clear set of practices that could be implemented for real-world use.

REFERENCE

1. iau, Keng, and Weiyu Wang. "Artificial Intelligence (AI) Ethics." *Journal of Database Management* 31, no. 2 (April 2020): 74–87.
2. , Zapata Flórez. "Cognitive Priority over Ethical Priority in Artificial Intelligence: The Primordial Philosophical Analysis in Artificial Intelligence." *Philosophy International Journal* 5, no. 4
3. R, Rishikesh. "Role of Ethics in Artificial Intelligence." *International Journal of Trend in Scientific Research and Development* Volume-2, Issue-5
4. Mohadi, Mawloud, and Yasser Tarshany. "Maqasid Al-Shari'ah and the Ethics of Artificial Intelligence." *Journal of Contemporary Maqasid Studies* 2,
5. Tasioulas, John. "Artificial Intelligence, Humanistic Ethics." *Daedalus* 151, no. 2 (2022): 232–43
6. Times of India, The Hindu ,India Toda and other referred journals