



Artificial intelligence for health care: open ethical challenges

Eliseo Sciarretta^a, Ferdinando Romano^{b*}, Edoardo Trebbi^b, Emilio Greco^a

^a *Link Campus University, Rome, Italy*

^b *Department of Public Health and Infectious Diseases, Sapienza University, Rome, Italy*

Abstract

Artificial intelligence applications are gaining traction in a variety of industries, with the promise of optimising resources and obtaining better outcomes through Big Data analysis. The potential of A.I. algorithms is generally recognised even in the healthcare field. However, healthcare, more than any other sector, must consider the ethical implications of employing these tools and what may happen if things do not go as planned. This essay addresses the key ethical dilemmas that are now open, establishing a classification based on functional elements and suggesting an order of significance to aid in their resolution.

Keywords: *artificial intelligence; health; ethics; acceptability; privacy; trust.*

1. Introduction

Artificial intelligence, defined as the ability to make computers perform tasks that require human-like intelligence (Garbuio and Lin 2019), or, in the words of one of its founders, John McCarthy, "the science of creating intelligent machines" (McCarthy, 2007), is finding application in an increasing number and diversity of fields.

Although the notion of artificial intelligence (A.I.) initially appeared in the mid-twentieth century (Hamet and Tremblay 2016; Russell and Norvig, 2016), it is only in the last 30 years that technology has advanced to the point where it is viable.

With today's computational power and connection, robots can learn from data, make autonomous decisions depending on external criteria, and interact with one another without human interference.

We frequently engage with A.I. applications in our daily lives without even realising it: ecommerce sites utilise these approaches to recommend new things for us to buy, much as streaming platforms do to tempt us to view new material that matches our preferences. Driver assistance technologies are becoming more common in automobiles, which may

*Corresponding author: Ferdinando Romano. E-mail address: ferdinando.romano@uniroma1.it.

eventually lead to self-driving vehicles. We have grown accustomed to using our voices to engage with smart speakers that offer us with information and services.

The possibilities are endless: in the military field, the use of drones with Artificial Intelligence allows actions to be carried out without involving human soldiers; in the industrial field, warehouse management can be automated and supplies established on the basis of endogenous and exogenous factors; in the government field, analysis algorithms combined with surveillance cameras are able to make predictions on possible crimes based on the behaviour of the subjects filmed; in the commercial field, as mentioned, it is possible to provide more personalised services based on the interests of users.

Artificial intelligence has already changed numerous economic sectors (Autor 2015), and in healthcare, A.I., which is regarded as the most promising technology for the sector (PwC 2017), is now perfectly integrated in numerous processes, focusing on their improvement and the optimisation of the necessary resources (Gutierrez 2020; Bellini et al., 2020), for example, for the analysis of clinical images or for monitoring people with chronic conditions.

Artificial intelligence is set to revolutionise healthcare by changing how professionals acquire data on patients as well as how they interact, while allowing machines to handle routine steps, flagging any discrepancies as needed, and unlocking the potential of Big Data to arrive at evidence-based clinical decisions (Chen and Decary 2020).

But, with so much at stake, particularly in the realm of health care, one has to consider how far one should go and at what cost.

Even Elon Musk, a visionary entrepreneur and proponent of new technology, has warned that the deployment of A.I. without defined regulations and profound knowledge might endanger civilization (Vincent 2017).

This paper seeks to systematise the ethical dangers involved in such an application of Artificial Intelligence, categorising them and assessing them in light of contemporary research on the issue.

2.Challenges detected

As just shown, the application of Artificial Intelligence in healthcare operations provides very exciting potential and will undoubtedly transform the way healthcare services are presented and delivered.

However, this will not happen in the near future because there are several open concerns that must be addressed.

To begin, while A.I. has made enormous achievements in recent years, it is still in its early stages of development, and will need to be developed and polished in the future in order to achieve its most ambitious goals.

But, beyond the technical issues that are not the focus of this contribution, it is necessary to draw attention to the acceptability of A.I.: we are increasingly dealing with algorithms that assist us in various activities, but are we willing to accept that this also occurs in the healthcare sector, where our health is at stake?

In other words, the application of A.I. to healthcare raises a slew of ethical concerns that must be addressed. The underlying premise of Artificial Intelligence is ethics (Jiang et al. 2021).

It is so important that the European Commission's Artificial Intelligence Expert Group (AI HLEG - High Level Expert Group on A.I.) deemed it necessary to define a set of requirements to ensure the trustworthiness of these systems, including: human

supervision, technical robustness, data security and privacy, transparency, equity, social welfare, and accountability (European Commission 2019).

Recently, the Organisation for Economic Co-operation and Development (OECD) compiled a similar list of ethical principles for A.I., reducing it to five: inclusiveness and well-being, human-centred values, transparency and explainability, security, and accountability (OECD 2021).

These recommendations, however, are intended for general uses of A.I., but their use in healthcare brings significant challenges.

We propose a subdivision of the ethical problems that may affect the impact of Artificial Intelligence in healthcare, consisting of 5 macro-categories, by removing the recommendations more related to the technical domain (robustness and security) and introducing concepts specifically related to healthcare that have emerged from the analysis of the literature recently produced in the field.

- Data protection and privacy (robustness is also included here)
- Accountability and government regulation
- Trust and comprehension (transparency and fairness, human-centred values)
- Doctor-patient relationship (also includes human supervision and social welfare)
- Trade-offs

However, a clarification is required: the importance of these difficulties is dependent on one's attitude toward A.I. As Kluge (2020) points out, there are two distinct viewpoints on the employment of these approaches in health care processes: the first regards A.I. as nothing more than a tool available to experts, to be employed under their rigorous supervision. The second, on the other hand, regards Artificial Intelligence as a type of creature that can function independently of human interference.

The concerns described in the first case are minor, because there is always a professional to take responsibility and assure the process's success. Artificial Intelligence is simply one tool at their disposal in this context, although it is a restricted application of A.I. However, in order to fully utilise the system's potential, it is required to move to the second mode, in which the system becomes autonomous, and it is here that the concerns described become crucial.

3. Data protection and privacy

The issue of data privacy, data transfer security, and probable fraudulent applications is perhaps the most debated in the literature, as well as the most felt by end users.

Data, in fact, is at the heart of how Artificial Intelligence works, which, in order to be effective, must examine a massive quantity of data, often seemingly unrelated, in search of correlations for the creation of algorithms. As a result, what we refer to as Big Data must be handled ethically (Boyd and Crawford 2012).

This applies equally to the broad application of Artificial Intelligence as it does to its application in healthcare, where data is by definition sensitive.

The concern that information about our health issues will become public or that it will be utilised for other reasons is a disincentive to the success of any digital system and thus constitutes the primary issue to be addressed.

Recently, the issue has come up in the discussion of a specific type of Artificial Intelligence application that has grown in popularity in the last five years: Intelligent Personal (or Virtual) Assistants, stand-alone systems capable of establishing a spoken dialogue with the user using advanced voice recognition and synthesis techniques to answer questions and provide additional services.

Studies on these systems are growing, with many emphasising the concerns of privacy and data security (Nallam et al. 2020; Liao et al. 2019). According to the findings of these investigations, purpose-built instruments capable of ensuring high security requirements are required for use in healthcare, but commercial solutions are not deemed totally dependable in this regard (Bickmore et al., 2018).

As a result, people who lack the skills to investigate privacy problems are more likely to rely on industry participants such as hospitals, clinics, or other healthcare providers rather than international IT businesses such as Apple or Google (Seifert, Christen and Martin, 2018).

This idea, which is applicable to IPAs, may be applied to all A.I. applications in healthcare. Especially since the use of Big Data in this subject might lead to disclosures that go beyond the purpose of the data collection. A.I. algorithms may be able to identify health concerns and create predictions about future circumstances by combining information taken from various visits, routine check-ups, consumption patterns, or socio-demographic data (Racine, Boehlen & Sample, 2019).

As a result, the utilisation of data and its security become even more important, necessitating stringent protocols.

Although the goal of this contribution is not to go into technical difficulties, it is worth noting that the robustness principles mentioned in the guidelines can easily fit into this category. Indeed, in order to maintain data security, it is required to act on the dependability and impenetrability of the hardware devices employed as well as the networks that transport this data (Jiang et al., 2021; O'Sullivan et al., 2019).

4.Accountability and government regulation

There is a need for clearly defined and controlled duties in the different service procedures in health care.

Our existing system assumes that duty belongs with the supplier of the service or care, but when errors occur, it is not always straightforward to establish the lines of accountability (Reddy, Fox and Purohit 2019).

This method is true if A.I. is seen as a simple tool because the practitioner is always held accountable. All practitioners in the sector have received particular and rigorous training, which includes adoption of an ethical code (Jiang et al. 2021).

But what if we take a step further and picture Artificial Intelligence making decisions on its own?

The regulatory structure is not prepared to sustain such an effect, and we would be faced with a serious medico-legal crisis.

Who is responsible for an intelligent agent's incorrect decision? (2016 Kingston) The engineer who created it, the technician who runs it, the doctor who recommended it, the patient who employs it, or the intelligent agent itself? (Jiang et al. 2021)

There is no attempt to govern such a scenario globally (Mitchell and Ploem 2018).

Before the entire evolution of A.I. in health to occur, these challenges must be resolved in order to have clear laws.

When an AI-based service is delivered, standards must be in place to establish which factors determine whether the service is supplied to certain patients and not others, and who is held accountable for its effectiveness (Fonseka, Bhat and Kennedy 2019).

It should also be emphasised that the scenario outlined thus far only looks at legal applications of these tools, with possibly honest mistakes. However, the scenario is

exacerbated further by the possibility of unlawful or criminal usage (King et al. 2020), which might be considerably more hazardous.

5.Trust and comprehension

If the issue of regulation is a prerequisite without which it is not even possible to begin experiments on an evolved concept of Artificial Intelligence, the eventual success of these techniques is inextricably linked to the level of transparency that they can provide to those who use them, and thus the level of trust that users are willing to place in them.

Artificial intelligence has already demonstrated that it can be more accurate than humans in clinical decision-making, particularly in fields such as radiology and dermatology (Goldhahn, Rampton, and Spinas 2018), but its development, in order to be fully effective, must adhere to human-centric principles throughout all stages of design (Kieslich, Keller and Starke 2022), in order to maximise interest and readiness for use (Berendt 2019; Jobin, Ienca and Vayena 2019).

To ensure that patients (and experts) can trust an A.I. algorithm's decisions, it is vital to ensure that they understand how the algorithm works and on what basis it makes specific decisions, so that they can constantly remain watchful and intuit if the decision deviates from the binary (Reddy, Fox and Purohit 2019).

However, by definition, the algorithms used take in massive amounts of data, far exceeding human processing capabilities, and their complexity makes them opaque, resulting in what is known as the “black box” problem (Amato et al. 2013; Shin and Park 2019), in which even experts in the field struggle to understand the motivation for some results (Basu, Engel-Wolf and Menzer 2020;).

Close collaboration between software developers and healthcare practitioners in the creation of algorithms is required to enhance this.

Greater transparency can also facilitate accountability (Diakopoulos 2016).

Even if comprehension is achieved, the issue of trust must still be addressed. Although the judgement made by an intelligent agent following rigid rules may appear to be the best option in certain circumstances, human experts have a better capacity to function on an ethical foundation, allowing them to make decisions in complicated scenarios that go beyond 'simple' data (Gelhaus 2011).

Furthermore, in order to trust, we must be absolutely convinced that no mistakes have occurred. In the realm of suicide prevention, for example, Artificial Intelligence tools are demonstrating their efficacy in large-scale screening to find at-risk persons in good time. However, before 'branding' a person as suicidal, one must be certain that one is not committing a grave error (Fonseka, Bhat and Kennedy 2019).

Decisions based on algorithm data, on the other hand, may not be totally right.

Several studies have revealed that Artificial Intelligence, like humans, is susceptible to biases since the learning phase, known as machine learning, is drawn from human expertise (Challen et al. 2019; Parikh, Obermeyer and Navathe 2019).

As a result, if the machine is trained to assign a certain weight to a parameter that is irrelevant, the problem is likely to persist without resolution, resulting in what is known as algorithmic discrimination, which is caused by faulty starting data or a flaw in the algorithm itself (Shin and Park 2019).

As a result, rather than providing equity and equal treatment, disparities are exacerbated (Rajkomar et al. 2018), particularly to the detriment of minorities (Kieslich, Keller and Starke 2022; Caliskan, Bryson and Narayanan 2017).

Biases in this context might be based on gender, sexual orientation, race, or social or economic reasons (Geis et al. 2019).

The 'confirmation bias,' in which more weight is given to data that confirm a certain explanation than to data that refute it, and the 'automation bias,' in which the practitioner blindly trusts the algorithm without carrying out the necessary checks to look for additional evidence, are typical biases that can occur in this context (Choudhury and Asan 2020).

Work must be done to reduce these biases, which are unlikely to be entirely removed.

Again, the situation outlined only considers concerns coming from genuine faults caused by problems with programmes that are inherently imperfect. These codes may clearly be improved with time, but in the medical industry, even little errors can be fatal (Jiang et al. 2021).

It should also be noted that data may be readily changed illegally to fit inside a specific model (Choudhury and Asan 2020).

6. Doctor-patient relationship

While the challenges discussed above are generally relevant to all uses of AI and are represented in the healthcare domain, with some exceptions, such as comprehensibility, the one concern how the doctor/patient relationship develops is extremely special and is strongly related to the issue of trust.

Participants in the previously mentioned Intelligent Personal Assistants research (Nallam et al. 2020) responded that while utilising IPAs to gather health information may be beneficial, they would feel safer if the information came from a medical expert.

Although the IPA may feature reminder, alert, and even monitoring functions, the preferred and most trustworthy source of information remains the physician, at least until approaches to increase the device's dependability and transparency are researched (Eiband et al. 2018).

If this is true for one of the most widely acknowledged commercial applications of Artificial Intelligence, it is reasonable to assume it to be true for others.

Returning to the example of suicide prevention, once a diagnostic of high risk is determined, this information must be presented to the patient by a professional, as the A.I. agent may lack empathy and make the situation worse (Fonseka, Bhat and Kennedy 2019).

In general, the doctor-patient relationship is a fiduciary covenant in which ethics plays a significant part (Kluge, 2020) to ensure that the patient understands that the doctor will always act in the patient's best interests. However, establishing a trust bond between a patient and artificial intelligence is significantly more difficult.

As Jiang et al. (2021) shown, when a doctor's diagnosis differs from that of an intelligent agent, the patient tends to believe physicians more, even if the intelligent agent is statistically less likely to be wrong.

7. Trade-offs

The last challenge category differs from the others in that it refers to the necessity to establish acceptable trade-offs when all other problems cannot be solved, or at least priority scales for their resolution.

Of course, perfection would be to be able to adequately answer all of the ethical principles, but in practise, this can cause many problems in the design of A.I. systems (Kieslich, Keller and Starke 2022), because some of the principles may conflict with each other, so the best compromise points must be identified (Binns and Gallo 2019; Köbis, Starke and Rahwan 2021).

According to research, these trade-off points, which are based on audience preferences, are affected by both the environment of use and individual traits (Pew Research Center 2018; Starke et al. 2021). We may readily accept the use of Artificial Intelligence to recommend which TV series to watch based on what we have already watched, but we are less likely to accept it if it has to advise us on which therapy to pursue.

According to a Pew Research Center (2018) research, there are four key worries about the use of A.I.: privacy violations, erroneous outcomes, the removal of the human component from critical choices, and A.I.'s inability to understand human complexity.

Further research will be required to properly understand which ethical values must be carefully adhered to and which ones individuals are ready to make exceptions to.

For example, privacy assurance is currently reported as a fundamental condition, but this is not true for all age groups of the public, and this perception is changing, not least because there is a lack of shared consensus on which data is to be considered private and which can be publicly accessible (Fonseka, Bhat and Kennedy 2019).

The most likely trade-off appears to be that intelligent agents will not be permitted to make autonomous judgments, at least in the short term, but will always be supervised by a human expert, especially in the context of healthcare (Jiang et al. 2021).

8. Conclusion

The analysis shows that in order for Artificial Intelligence to be used effectively in health care processes, certain challenges, primarily concerning ethical aspects, must be addressed and broadly acceptable solutions must be found, as well as waiting for the technology to mature so that it can fulfil its initial promises.

The issues raised throughout the debate are all essential, but as previously said, one of the requirements is to establish priority levels in order to reach compromises.

The findings of the investigation led to the conclusion, as supported by Kieslich, Keller, and Starke (2022), that the most pressing issue is one of regulations. Accountability cannot be established without a specific and globally agreed-upon regulatory framework, hence this must be regarded a fundamental prerequisite for the growth of A.I.

The European Union recently took a step in this direction by proposing a legislative framework for the use of A.I. (European Commission 2021), which is based on a risk categorization of specific uses of the technology.

Instead, among those explored, the topic of algorithm comprehensibility might be placed in the background for the time being. Although it is clear that a higher level of comprehensibility increases trust in A.I., and thus its acceptability to the public, it may be sufficient for the time being to focus on the effectiveness of the results, showing people that, for example, a decision made by a given algorithm has higher rates of correctness than one made by a professional who cannot rely on the same amount of data.

The effectiveness of the system could thus be a sufficient picklock to overcome public reluctance, at least in situations where the perceived risk is not so high, so that more time could be given to improving the comprehensibility of these systems from the standpoint of eXplainable Artificial Intelligence.

The other two difficulties, those concerning privacy on the one hand and the doctor/patient relationship on the other, might be viewed as intermediate, but with adequate differentiation.

Data security, like liability, must be addressed before we can begin applying Artificial Intelligence techniques on a big scale, particularly by delving into patients disaggregated personal data.

The doctor-patient relationship, on the other hand, is by definition a far more subjective topic, based only on one's own viewpoint, and hence not necessarily valid for everyone, making tackling this problem less essential.

Finally, it is necessary to reiterate the clarification made at the beginning of the contribution, namely that the problems discussed take on a much greater importance when one refers to the use of Artificial Intelligence with a high degree of decision-making autonomy, which is currently hardly feasible both technically and due to the problems described. For these reasons, it is advisable to utilise A.I. in healthcare under expert supervision.

References

- Amato, F., Lopez, A., Pena-Mendez, E.M., Vanhara, P., Hampl, A. & Havel, J. (2013). "Artificial neural networks in medical diagnosis". *J Appl Biomed* 2013; 11: 47–58.
- Autor, D.H. (2015). "Why are there still so many jobs? The history and future of workplace automation". *J Economic Perspectives*. 2015; 29:3-30.
- Basu, T., Engel-Wolf, S., & Menzer, O. (2020). "The ethics of machine learning in medical sciences: Where do we stand today?" *Indian Journal of Dermatology*, 65(5), 358.
- Bellini, V., Guzzon, M., Bigliardi, B., Mordonini, M., Filippelli, S., & Bignami, E. (2020). "Artificial intelligence: a new tool in operating room management. Role of machine learning models in operating room optimization". *Journal of medical systems*, 44(1), 1-10.
- Berendt, B. (2019). "AI for the common good?! Pitfalls, challenges, and ethics pen-testing". *Paladyn, Journal of Behavioral Robotics* 10(1): 44–65.
- Bickmore, T. W., Trinh, H., Olafsson, S., O'Leary, T. K., Asadi, R., Rickles, N. M., & Cruz, R. (2018). "Patient and consumer safety risks when using conversational assistants for medical information: An observational study of Siri, Alexa, and Google Assistant". *Journal of Medical Internet Research*, 20(9), e11510. <https://doi.org/10.2196/11510>.
- Binns, R., & Gallo, V. (2019). "Trade-offs". Available: <https://ico.org.uk/about-the-ico/news-and-events/ai-blog-trade-offs/>
- Boyd, D., & Crawford, K. (2012). "Critical questions for Big Data: Provocations for a cultural, technological, and scholarly phenomenon". *Information, Communication & Society*. 15(5):662–679.
- Caliskan, A., Bryson, J.J. & Narayanan, A. (2017). "Semantics derived automatically from language corpora contain human-like biases". *Science* 2017; 356: 183–186.
- Challen, R., Denny, J., Pitt, M., Gompels, L., Edwards, T., & Tsaneva-Atanasova, K. (2019). "Artificial intelligence, bias and clinical safety". *BMJ Quality & Safety*, 28(3), 231-237.
- Chen, M., & Decary, M. (2020). "Artificial intelligence in healthcare: An essential guide for health leaders". In *Healthcare management forum* (Vol. 33, No. 1, pp. 10-18). Sage.
- Choudhury, A., & Asan, O. (2020). "Human factors: bridging artificial intelligence and patient safety". In *Proceedings of the International Symposium on Human Factors and Ergonomics in Health Care* (Vol. 9, No. 1, pp. 211-215). Sage.
- Diakopoulos, N. (2016). "Accountability in algorithmic decision making". *Communications of the ACM* 59(2): 56–62.
- Eiband, M., Schneider, H., Bilandzic, M., Fazekas-Con, J., Haug, M., & Hussmann, H. (2018). "Bringing transparency design into practice". *23rd International Conference on Intelligent User Interfaces* (pp. 211–223). <https://doi.org/10.1145/3172944.3172961>.

European Commission (2019). "Ethics guidelines for trustworthy AI". Available: <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>

European Commission (2021). "Proposal for a regulation of the European Parliament and of the council: Laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain union legislative acts". Available: https://ec.europa.eu/newsroom/dae/document.cfm?doc_id=75788

Fonseka, T.M., Bhat, V., & Kennedy, S.H. (2019). "The utility of artificial intelligence in suicide risk prediction and the management of suicidal behaviors". *Australian & New Zealand Journal of Psychiatry*, 53(10), 954-964.

Garbuio, M., & Lin, N. (2019). Artificial intelligence as a growth engine for health care startups: Emerging business models. *California Management Review*, 61(2), 59-83.

Geis, J. R., Brady, A. P., Wu, C. C., Spencer, J., Ranschaert, E., Jaremko, J. L., ... & Kohli, M. (2019). "Ethics of artificial intelligence in radiology: summary of the joint European and North American multisociety statement". *Canadian Association of Radiologists Journal*, 70(4), 329-334.

Gelhaus, P. (2011). "Robot decisions: on the importance of virtuous judgment in clinical decision making". *Journal of Evaluation in Clinical Practice*, 17(5), 883-887.

Goldhahn, J., Rampton, V., & Spinaz, G. A. (2018). "Could artificial intelligence make doctors obsolete?" *Bmj*, 363.

Gutierrez, G. (2020). "Artificial intelligence in the intensive care unit". *Crit Care* 2020; 24: 101.

Hamet, P., & Tremblay, J. (2016). "Artificial intelligence in medicine". *Metabolism* 2017; 69s: S36-S40.

Jiang, L., Wu, Z., Xu, X., Zhan, Y., Jin, X., Wang, L., & Qiu, Y. (2021). "Opportunities and challenges of artificial intelligence in the medical field: Current application, emerging problems, and problem-solving strategies". *Journal of International Medical Research*, 49(3), 03000605211000157.

Jobin, A., Ienca, M., & Vayena, E. (2019). "The global landscape of AI ethics guidelines". *Nature Machine Intelligence* 1(9): 389-399.

Kieslich, K., Keller, B., & Starke, C. (2022). "Artificial intelligence ethics by design. Evaluating public perception on the importance of ethical design principles of artificial intelligence". *Big Data & Society*, 9(1), 20539517221092956.

King, T.C., Aggarwal, N., Taddeo, M., & Floridi, L. (2020). "Artificial intelligence crime: An interdisciplinary analysis of foreseeable threats and solutions". *Science and engineering ethics*, 26(1), 89-120.

Kingston, J. (2016). "Artificial Intelligence and Legal Liability". *Research and Development in Intelligent Systems XXXIII*. Springer, Cham.

Kluge, E.H.W. (2020). "Artificial intelligence in healthcare: Ethical considerations". In *Healthcare management forum* (Vol. 33, No. 1, pp. 47-49). Sage.

Köbis, N., Starke, C., & Rahwan, I. (2021). "Artificial intelligence as an anti-corruption tool (AI-ACT): Potentials and pitfalls for top-down and bottom-up approaches". <https://arxiv.org/abs/2102.11567>.

Liao, Y., Vitak, J., Kumar, P., Zimmer, M., & Kritikos, K. (2019). "Understanding the role of privacy and trust in intelligent personal assistant adoption". In N. G. Taylor, C. Christian-Lamb, M. H. Martin, & B. Nardi (Eds.), *Information in contemporary society* (pp. 102-113). Springer International Publishing.

McCarthy, J. (2007). "What is artificial intelligence?" Available: <http://www-formal.stanford.edu/jmc/whatisai/>.

Mitchell, C., & Ploem, C. (2018). "Legal challenges for the implementation of advanced clinical digital decision support systems in Europe". *J Clin Transl Res* 2018; 3: 424–430.

Nallam, P., Bhandari, S., Sanders, J., & Martin-Hammond, A. (2020). "A question of access: Exploring the perceived benefits and barriers of intelligent voice assistants for improving access to consumer health resources among low-income older adults". *Gerontology and Geriatric Medicine*, 6, 2333721420985975.

O'Sullivan, S., Nevejans, N., Allen, C., Blyth, A., Leonard, S., Pagallo, U., Holzinger, K., Holzinger, A. Sajid, M.I., & Ashrafian, H. (2019). "Legal, regulatory, and ethical frameworks for development of standards in artificial intelligence (AI) and autonomous robotic surgery". *The international journal of medical robotics and computer assisted surgery*, 15(1), e1968.

OECD (2021). "Recommendation of the Council on Artificial Intelligence". Available: <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>

Parikh, R.B., Obermeyer, Z., & Navathe, A.S. (2019). "Regulation of predictive analytics in medicine". *Science*, 363(6429), 810-812.

Pew Research Center (2018). "Public attitudes towards computer algorithms". Available: <https://www.pewresearch.org/internet/2018/11/16/public-attitudes-toward-computer-algorithms/>.

PwC (2017). "What doctor? Why AI and robotics will define New Health". Available: <https://www.pwc.com/gx/en/industries/healthcare/publications/ai-robotics-new-health/ai-robotics-new-health.pdf>

Racine, E., Boehlen, W., & Sample, M. (2019). "Healthcare uses of artificial intelligence: Challenges and opportunities for growth". In *Healthcare management forum* (Vol. 32, No. 5, pp. 272-275). Sage.

Rajkomar, A., Hardt, M., Howell, M. D., Corrado, G., & Chin, M. H. (2018). "Ensuring fairness in machine learning to advance health equity". *Annals of internal medicine*, 169(12), 866-872.

Reddy, S., Fox, J., & Purohit, M. P. (2019). "Artificial intelligence-enabled healthcare delivery". *Journal of the Royal Society of Medicine*, 112(1), 22-28.

Russell, S., & Norvig, P. (2016). "Artificial Intelligence: A Modern Approach". 3rd ed. Essex: Pearson Education Limited.

Seifert, A., Christen, M., & Martin, M. (2018). "Willingness of older adults to share mobile health data with researchers". *GeroPsych: The Journal of Gerontopsychology and Geriatric Psychiatry*, 31(1), 41–49. <https://doi.org/10.1024/1662-9647/a000181>.

Shin, D. & Park, Y.J. (2019). "Role of fairness, accountability, and transparency in algorithmic affordance". *Computers in Human Behavior* 98: 277–284.

Starke, C., Baleis, J., Keller, B., & Marcinkowski, F. (2021). "Fairness perceptions of algorithmic decision-making: A systematic review of the empirical literature". *arXiv:2103.12016*. <http://arxiv.org/pdf/2103.12016v1>.

Vincent, J. (2017). "Elon musk says we need to regulate AI before it becomes a danger to humanity". *The Verge*. Available: <https://www.theverge.com/2017/7/17/15980954/elon-musk-ai-regulation-existential-threat>.

Received 08 August 2022, accepted 02 December 2022