

Concept Paper

Not peer-reviewed version

Analysis of Artificial Intelligence in the Post-Pandemic Era

[Felipe Leal Valentim](#) *

Posted Date: 28 January 2026

doi: 10.20944/preprints202601.2196.v1

Keywords: artificial intelligence; ethics; post-pandemic



Preprints.org is a free multidisciplinary platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This open access article is published under a [Creative Commons CC BY 4.0 license](#), which permit the free download, distribution, and reuse, provided that the author and preprint are cited in any reuse.

Disclaimer/Publisher's Note: The statements, opinions, and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions, or products referred to in the content.

Concept Paper

Analysis of Artificial Intelligence in the Post-Pandemic Era

Felipe Leal Valentim

Independent Researcher, Brazil; felipelealvalentim@usp.br

Abstract

During the pandemics, the positive value of technologies was emphasized. In the post-pandemic era, shortly after the easing of confinement, the negative values were also re-evidenced. Despite the noted depreciation, it is agreed that technological advancement will always have a positive balance, but that there cannot be injustice due to lack of access to technology. This can be the subject of studies on digital inclusion. In turn, the set of values and practices that seek to ensure that the development and use of artificial intelligence (AI) systems are safe, fair, and responsible is discussed in the ethical and moral sciences of AI. This work presents a write-up as an attempt to generalize the framework presented in the work done by Michalski et al. (2025) and discuss a) norms for evaluating needs and areas of application, b) definition of values of the methods, and c) definition of criteria for comparing techniques.

Keywords: artificial intelligence. ethics. post-pandemic

RESUMO: Durante a pandemia o valor positivo das tecnologias foi enfatizado. Na era pós-pandemia, logo após o desconfinamento, os desvalores também foram re-evidenciados. Apesar da depreciação notada, é concordado que o avanço tecnológico vai sempre ter um balanço positivo, mas que não pode haver injustiça devido à falta do acesso à tecnologia. Isso pode ser pauta de estudos de inclusão digital. Por sua vez, o conjunto de valores e práticas que buscam garantir que o desenvolvimento e o uso de sistemas de inteligência artificial (IA) sejam seguros, justos e responsáveis é discutido nas ciências éticas e morais de IA. Esse trabalho apresenta uma resenha como tentativa de generalizar o *framework* do trabalho feito por Michalski et al., (2025) e discutir a) normas para avaliar necessidades e áreas de aplicações, b) definição de valores dos métodos, e c) definição de critérios para comparação das técnicas.

PALAVRAS-CHAVE: inteligência artificial. ética. pós-pandemia.

1. Introduction

During the pandemic, the positive value of technologies was emphasized in applications such as autonomous driving, medical assistance, media, finance, industrial robots, domestic robots, and internet services (Huang et al., 2023). Artificial intelligence allowed people to stay close to loved ones, remain healthy at home, or work in home office with increased profits for both employer and employee. In the post-pandemic era, the situation regressed, and the negative aspects were re-evidenced. The end of the sudden lockdown revealed dystopias, unmet expectations, and naivety regarding promises of happiness and profit through the use of artificial intelligence. This perhaps because the market reimposed norms on software production that do not add value to products but supposedly increase indirect profitability due to policy and marketing. In other words, it lost some of its balance between value and competitiveness.

2. Methodology

This review explores how to define ethical values and how to transfer these values to technology, especially artificial intelligence. Due to influences, crimes, security concerns, etc., discussions about artificial intelligence have moved beyond academia and entered into the political, legal, and religious realms. Since the goal is to assist or automate decision-making, reintroducing moral values into artificial intelligence is desirable; in other words, implementing moral agency. Initially, moral agency (Etzioni et al., 2017) can be planned focusing on tools that prevent deaths, bring people together at home, save money, improve quality of life at work, increase longevity, and so on. Technically, it is possible to discuss the implementation of moral agency in the development steps to transfer human values to the applications, such as interpretability, applicability, profitability, usability, innovation, originality, and so on, along with other ethical values.

Transferring human intelligence to AI lies also in incorporating ethical values during the planning and development of technologies. For example, interpretability. Solving mathematical equations is easy for a computer, even without knowledge of some variables. However, for humans, interpreting the results remains difficult. Another example of an ethical value is applicability, which characterizes in what terms it is worthwhile to apply the technology, such as monitoring, diagnosis, alarms, and so on. Furthermore, there is profitability, that is, interest in products that help accumulate wealth at the expense of superfluous software; or usability, which considers the interests and capabilities of users' minds. Defining ethical criteria and values for choosing a method is a task in itself (Michalski et al., 2025). Figure 1 shows an example of a framework for transferring ethical-moral values to artificial intelligence.

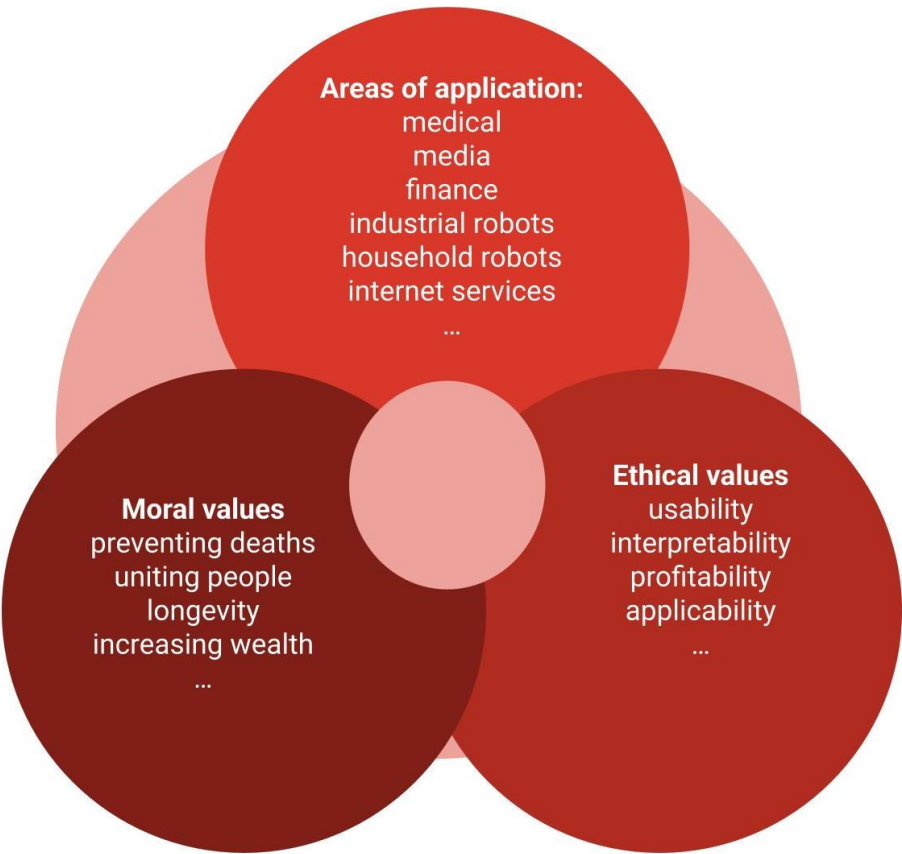


Figure 1. Examples of frameworks for transferring ethical and moral values to artificial intelligence.

Human moral essence seems to originate from the use of tools, as well as in reading and writing. However, it is only seen in interactions among humans between humans and other living beings. Computer tools and language are different from those of humans, so in essence, the morality of robots

and human are different. What can be done, therefore, is to consider human morality and ethics in the development and in the use of technology.

In computer science, ethical values are defined in layers. The first layer, data, defines criteria such as reproducibility, data consistency, coherence of hypotheses about the data, and so on. The second layer, the algorithm layer, assigns values to the techniques themselves, such as monitoring, prediction, diagnosis, etc. The third layer, the algorithm's interface with the user, encompasses the human values that the algorithm can assimilate. That is, which ethical values can be translated into numbers. Examples include interpretability, usability, complexity, and so on. Assigning values to techniques allows human intelligence to be applied to the use or development of the technology.

One way to define ethical criteria for the techniques in machine learning is through literature reviews, as done in Michalski et al. (2025). The proposed review method associates the terms of the values, the application label, and the names of the techniques when these appear in the same study. From there, the frequency of these associations is analysed. It is then possible to quantify whether the "random forest" machine learning method, for example, has high usability. That is, whether "random forest" appears in the same studies where keywords that define usability values appear.

When dealing with specific application area, such as engineering, medicine, etc., the most values are defined by interest. "Diagnosing," "predicting," "clustering," "recognition," "monitoring," among others, are terms that can be associated with the values of these applications. Furthermore, various disciplines offer techniques for solving artificial intelligence problems, such as statistics, mathematics, linguistics, and, most naturally, computer science.

Conclusions

The form of human intelligence is conditioned by the fact that we have senses and a human brain capable of performing calculations and executing tasks. Artificial intelligence is based on a) imitating the reasoning of experts and professionals, b) learning how the human mind performs calculations and tasks, or c) the biological structures that allow the brain to use human intelligence. There is also the case of 4) obtaining inspiration from nature. For example, there is interest in imitating medical reasoning when calculating the probabilities of a diagnosis given test results. Since the training that enables the doctor's mind to perform the task is standardized in medicine, it is possible to learn the reasoning pattern and transform it into an algorithm. With Bayesian networks (VALENTIM, F. L, 2007) it is intuitive to model the medical reasoning that performs this task. Another example is the engineer capable of making predictions from observations of temporal data; for this, there are numerous forecasting algorithms (e.g., Hyndman, R. J., & Khandakar, Y. 2008).

In addition to mimicking the reasoning of experts, there is also the ability to understand how the human mind performs everyday tasks such as classifying. Furthermore, there is the form of intelligence conditioned by situations. Namely, objective, subjective, descriptive reasoning, and others that depend on environments and cultures. There can be objective and subjective solutions to the same problem. In this case, heuristics are attempts to capture and understand how the human mind solves everyday problems. Finally, there is the branch of artificial intelligence that was inspired by the biological structure of the human brain, with neurons, synapses, and biological circuits. This branch was responsible for proposing techniques such as "neural networks".

Authors' contribution statement: Single author.

Declaration of availability of research data: All the data supporting the results of this study were published in the article itself.

Declaration of conflict of interest: The authors declare that there is no conflict of interest.

References

1. HUANG, Changwu; Zeqi, ZHANG; MAO, Bifei; YAO, Xin. An Overview of Artificial Intelligence Ethics. IEEE TRANSACTIONS ON ARTIFICIAL INTELLIGENCE, v. 4, n. 4, 2023.

2. ETZIONI, Amitai; ETZIONI, Oren. Incorporating Ethics into Artificial Intelligence. J Ethics, v. 21, 2017.
3. MICHALSKI, Miguel A. De C.; MURAD, Carlos A., KAHIWAGI, Fabio N., De SOUZA, Gilberto F. M.; Da SILVA, Halley J. B.; CÔRTEZ, Hyghor M. A Multi-Criteria Framework for Selecting Machine Learning Techniques for Industrial Fault Prognosis. IEEE. v: 13, 2025.
4. VALENTIM, F. L. . Estudo e Implementação de Algoritmos de Inferência e Aprendizado em Redes Bayesianas. Monografia, UFLA, 2007.
5. HYNDMAN, R. J., & KHANDAKAR, Y. (2008). Automatic Time Series Forecasting: The forecast Package for R. Journal of Statistical Software, 27(3), 1–22.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.