

Rethinking the scientific method in the age of artificial intelligence: a philosophy of science perspective

Madaniyah^{1*}, Ahmad Syukri², Mohd. Arifullah²

¹SDN 001 Keramat, Kepulauan Anambas, Kepulauan Riau, Indonesia

²Universitas Islam Negeri Sulthan Thaha Saifuddin Jambi, Indonesia

*Correspondence author: madaniyah.94@admin.sd.belajar.id

DOI: <https://doi.org/10.65881/innovation.v1i2.95>

ARTICLE INFO

History:

Received: 06-08-2026

Revised: 06-12-2026

Accepted: 06-13-2026

Published: 06-17-2026

Keywords:

artificial intelligence;
scientific method;
philosophy of science;
epistemology;
data ontology.

ABSTRACT

Purpose: to analyze the transformation of the scientific method in the era of Artificial Intelligence (AI) through a philosophical framework, focusing on ontological, epistemological, and axiological changes, as well as the emergence of human-machine integration in knowledge production.

Method: this study uses a qualitative descriptive design with a library research approach. Data were drawn from ten international journal articles (2020–2025) selected purposively and analyzed using qualitative content and thematic analysis to examine ontological, epistemological, and axiological changes in the scientific method in the context of Artificial Intelligence.

Findings: AI reshapes the scientific method by shifting its ontology toward data-driven reality, its epistemology from causal explanation to algorithmic prediction, and its axiology toward ethical concerns such as bias and fairness. Scientific practice is increasingly characterized by human-machine collaboration, combining AI's computational power with human interpretive judgment.

Implications: AI reshapes scientific foundations and requires rethinking objectivity, explanation, and ethics, while strengthening the role of philosophical reflection to ensure transparent and responsible research.

Originality: lies in its integration of ontological, epistemological, and axiological perspectives to show AI as a driver of transformation in the scientific method.



Open access article under CC-BY-SA license.



Introduction

The development of Artificial Intelligence (AI) over the past two decades has brought fundamental changes to the way scientific knowledge is produced, verified, and utilized (Wang et al., 2023). While in the classical tradition of the scientific method,

researchers played a dominant role in formulating problems, formulating hypotheses, collecting data, analyzing findings, and drawing generalizations, in the era of intelligent computing, some of these stages are now carried out through collaboration between humans and algorithmic systems that can learn from data (Ren et al., 2023; Zhu et al., 2023). AI no longer functions simply as a statistical tool; it has evolved into a new epistemic infrastructure that shapes how knowledge is constructed (Alvarado, 2023). Through machine learning, deep learning, natural language processing, and generative AI, AI systems can discover patterns that humans cannot easily recognize, predict complex phenomena, and automatically synthesize information (Razzaq & Shah, 2025). This transformation demonstrates that scientific knowledge is no longer solely the product of human rationality but also arises from the dynamic interaction among humans, data, and machines (Tsvetkova et al., 2024).

The philosophy of science is regaining relevance as fundamental questions about the nature of knowledge, the validity of truth, and the purposes for which it is used become increasingly important. Floridi & Chiriatti (2020) assert that modern AI systems possess impressive generative capabilities but are not identical to the rational understanding humans possess. Therefore, practically useful AI output still requires epistemological justification to be accepted as valid scientific knowledge. Furthermore, Dwivedi et al. (2021) show that AI has evolved into a multidisciplinary phenomenon that broadly impacts governance, economics, education, and scientific research. At the same time, the development of large-scale language models demonstrates AI's ability to generate synthetic text and knowledge, but also raises the risk of bias, hallucinations, and the reproduction of inequalities derived from training data (Bender et al., 2021).

These developments present both new phenomena and fundamental problems in scientific research practices. One key issue is the tension between technical effectiveness and scientific legitimacy. Many AI models demonstrate very high predictive performance, but the processes that produce these decisions are often difficult to understand and explain (Hassija et al., 2024). This situation is known as the black-box problem: a system produces accurate output, but the underlying inference path is not transparent. In the scientific tradition, scientific truth is measured not only by the accuracy of results but also by the ability to explain, test, and replicate the processes that produce those findings. Therefore, high predictive accuracy does not automatically guarantee the epistemological validity of an AI system. Eschenbach (2021) asserts that the primary source of distrust in AI lies not solely in the possibility of output errors, but in the lack of clarity about the reasoning behind machine-generated decisions.

This phenomenon is evident in the healthcare sector, where AI is used to detect diseases, predict patient risk, and support clinical diagnosis. Despite their high accuracy, AI-generated recommendations are often questioned when doctors cannot explain the rationale for the system's decision-making (Topol, 2019). In Europe, the discourse on AI focuses more on issues of regulation, fairness, and protecting citizens' rights from algorithmic bias. Mittelstadt (2019) emphasizes that ethical principles alone are insufficient to guarantee fairness in AI without concrete implementation mechanisms. Meanwhile, in many developing countries, challenges arise related to data quality, limited digital infrastructure, and low AI literacy among researchers. These differing contexts demonstrate that AI's problems are not universal, but rather influenced by varying social, economic, and institutional conditions.

Various solutions have been proposed to address these challenges, including Explainable AI (XAI), human-in-the-loop approaches, algorithmic audits, and ethical governance frameworks. Rai (2020) explains that the transformation from black box to

glass box is an urgent need for AI systems to be accepted in scientific practice and public decision-making. On the other hand, research by Brynjolfsson et al. (2023) shows that implementing generative AI can significantly increase productivity, though it still requires human supervision to maintain the quality and accountability of the results. However, most of the proposed solutions still focus on technical and operational aspects. However, the root of the problem lies in a more fundamental dimension: how science defines truth, objectivity, scientific explanation, and responsibility in the AI era. Therefore, a philosophy of science approach is crucial for understanding the transformations occurring in the methodological foundations of contemporary research.

The study of AI is indeed developing rapidly, but most research still positions the philosophy of science as a normative background, rather than as a primary analytical tool. Studies focus more on improving model accuracy, computational efficiency, mitigating data bias, or the social impacts of AI use (Kumar et al., 2026; Shah & Sureja, 2025; Sreerama & Krishnamoorthy, 2022). Consequently, discussions on how AI is changing the fundamental structure of the scientific method remain relatively limited. However, the advent of AI is challenging several classic assumptions in science. First, the assumption that researchers are the sole primary subjects in knowledge production has become increasingly problematic as algorithms now play a role in discovering patterns and building inferences. Second, the assumption that causal explanation is the primary goal of science has been questioned, as many AI models can produce highly accurate predictions without providing adequate explanations. Third, the assumption of objectivity as a value-free condition becomes increasingly complex because data, model design, and training parameters always involve human decisions that are not entirely neutral (Bender et al., 2021).

Thus, there is a research gap in the absence of studies that systematically integrate ontological, epistemological, and axiological perspectives to explain the paradigm shift of scientific methods in the AI era. The novelty of this research lies in developing a conceptual framework that positions these three main pillars of the philosophy of science as an analytical foundation for understanding the transformation of AI-based scientific methods. Ontologically, this research examines the status of data, digital reality, and computational representation as new objects of science. Epistemologically, this research evaluates whether algorithmic predictions can be accepted as scientific knowledge without adequate causal explanations (Floridi & Chiriatti, 2020). Axiologically, this research examines how the development and utilization of AI should be directed towards social justice, public benefit, and sustainability (Mittelstadt, 2019). This research aims to analyze the contribution of the philosophy of science to the formulation of a new paradigm for scientific methods in AI-based research. Specifically, this research seeks to identify fundamental changes in the stages of the scientific method when AI is involved in problem formulation, data processing, inference, and decision-making. explain the ontological, epistemological, and axiological problems that arise from the integration of AI into research practice; and formulate a conceptual model of the scientific method that is adaptive to technological developments without ignoring the principles of rationality, transparency, accountability, and humanity.

This research is crucial because the development of AI not only brings technological innovation but also raises philosophical challenges that directly impact the validity and legitimacy of science. Without an adequate philosophical foundation, the use of AI has the potential to shift the orientation of science from understanding to mere prediction, from explanation to efficiency, and from public responsibility to

technological domination. Therefore, this research is expected to make a theoretical contribution by strengthening the study of the philosophy of science in understanding the relationships among humans, data, and machines in the production of scientific knowledge. Furthermore, this research is expected to provide practical contributions to researchers, academics, and policymakers by helping them design AI-based research practices that are more transparent, accountable, and oriented towards the public interest. Thus, this article offers a more comprehensive philosophical framework for understanding the transformation of scientific methods in the AI era while providing direction for a science that remains critical, reflective, and human-centered.

Method

This research uses a descriptive qualitative approach with a library research method. The qualitative approach was chosen because the research objective is not to test statistical hypotheses or to quantitatively measure relationships between variables, but rather to deeply understand the transformation of scientific methods in the era of Artificial Intelligence (AI) and the contribution of the philosophy of science to explaining these changes. In the qualitative paradigm, reality is understood as a construction of meaning that must be interpreted through a critical examination of the ideas, texts, and contexts that underlie it. Therefore, a library study is a relevant approach because the primary objects of research are theories, concepts, philosophical arguments, and published research results. According to Aspers & Corte (2021), qualitative research is oriented towards a deep understanding of social and conceptual meanings, allowing researchers to explore complex phenomena through interpretations of various sources of knowledge. In the context of this research, the library study is used to examine the development of thinking on AI, changes in scientific method procedures, and the ontological, epistemological, and axiological implications arising from the integration of intelligent technology into modern research practices.

The research design is descriptive-analytical. The descriptive aspect is used to map the development of ideas, issues, and challenges related to AI-based research in the current scientific literature. Meanwhile, the analytical aspect is used to critically examine, compare, and synthesize various academic perspectives to produce a more comprehensive conceptual construct. This approach allows the research not only to present a summary of existing literature but also to develop theoretical arguments about the new paradigm of scientific methods in the AI era. The research process is carried out in stages, including identifying the study's focus, conducting a literature review, selecting data sources, categorizing findings by theme, and critically interpreting the findings to produce a conceptual synthesis.

The research data sources consist of scientific documents relevant to the themes of philosophy of science, scientific methods, and Artificial Intelligence. In the literature review, scientific documents serve as the research subject because they are the primary source of information analyzed. The primary data for this study consists of ten international scientific articles published between 2020 and 2025. This period was selected because the development of AI, particularly generative AI and large-scale language models, has accelerated significantly in recent years, resulting in increasingly intensive philosophical and methodological discourse. The number of articles was limited to maintain the depth of analysis and to allow for a more comprehensive interpretation of each source.

Data sources were selected using purposive sampling, a deliberate selection of documents based on specific criteria relevant to the research objectives. Inclusion

criteria included: (1) articles discussing AI in the context of scientific research, methodology, ethics, or philosophy of science; (2) published in reputable international journals and indexed in credible academic databases; (3) fully available through academic platforms such as Google Scholar, Scopus, Springer, Elsevier, Wiley, or the ACM Digital Library; and (4) clearly making a conceptual contribution to the discussion of AI-based scientific method transformations. Conversely, popular articles, non-academic opinion pieces, manuscripts that have not undergone peer review, and documents not directly related to the research focus were excluded from the analysis. The use of purposive sampling in qualitative research is appropriate because the research objective is not to make statistical generalizations, but rather to gain an in-depth understanding from information-rich, conceptually relevant sources (Campbell et al., 2020).

The primary research instrument is the researcher (human instrument), who plays a role in determining data sources, conducting critical readings, interpreting meaning, developing analytical categories, and compiling a conceptual synthesis. In qualitative research, the researcher is the central instrument because the process of interpreting data relies heavily on the researcher's reflective and interpretive abilities (Aspers & Corte, 2021). To improve the systematic nature of data collection, a literature analysis matrix was used, including the article's identity, research objectives, methodological approach, key findings, the philosophical concepts of science discussed, AI issues studied, and their relevance to changes in scientific methods. The matrix facilitated the coding process, comparison of findings between articles, and the preparation of a theoretical synthesis. Instrument validity was maintained through alignment of analysis indicators with research objectives, while consistency in data recording and grouping ensured the reliability of the analysis process.

The research procedure was carried out through several systematic stages. The first stage was the formulation of the study focus and research questions based on the issue of the transformation of scientific methods in the AI era. The second stage involved a literature search using keywords such as philosophy of science, scientific method, artificial intelligence, AI ethics, and explainable AI. The third stage was the selection of articles based on predetermined inclusion and exclusion criteria. The fourth stage involved a close reading of the selected articles, followed by recording the data in an analysis matrix. The fifth stage involved grouping the data into main themes covering ontology, epistemology, axiology, and the transformation of scientific methods. The final stage was a critical interpretation and conceptual synthesis that connected the various literature findings into a coherent framework. This procedure enabled the literature review to produce systematic and academically accountable findings (Xiao & Watson, 2019).

Data analysis was conducted using a combination of qualitative content analysis and thematic analysis. Content analysis was used to identify key concepts, argumentation patterns, and thought trends emerging in the analyzed documents. Meanwhile, thematic analysis was used to organize the data into themes relevant to the research focus. Braun & Clarke (2021) explain that thematic analysis is an effective method for systematically and flexibly identifying, analyzing, and interpreting patterns of meaning in textual data. The analysis began with data reduction, sorting relevant information from each article. Next, the data were coded regarding ontological aspects (the nature of data and digital reality), epistemological (validity, objectivity, justification, reproducibility, and explainability), axiological (ethics, justice, social utility, and public responsibility), and changes in scientific methods in the AI era. The coding results were then grouped into broader thematic categories to facilitate interpretation and synthesis.

To enhance the credibility of the research findings, source triangulation techniques were employed to compare perspectives, arguments, and findings across various articles, journals, and disciplines. Furthermore, coding and interpretation processes were repeated to ensure consistency and accuracy of the resulting meanings. Through this approach, this research not only produced a literature summary but also developed a conceptual model of a new paradigm for AI-based scientific methods rooted in the principles of the philosophy of science. Thus, data analysis served as a bridge between the literature findings and the theoretical contributions of this research (Nowell et al., 2017).

Results and discussion

Based on an analysis of ten international articles selected through purposive sampling, it was found that the transformation of scientific methods in the Artificial Intelligence (AI) era has not only occurred at the level of technology use as a research tool, but has also affected the fundamental structure of scientific knowledge production. The results of the coding process, data reduction, and thematic analysis identified four main themes that consistently emerged in the literature, namely: (1) ontological shifts in the objects of science, (2) epistemological transformations in the process of knowledge validation, (3) axiological challenges related to ethics and social responsibility, and (4) the emergence of a human-machine integration model as a new paradigm for the scientific method. These four themes indicate that AI has changed the traditional relationship between researchers, data, methods, and research results. While in the classical scientific paradigm researchers have full control over the entire research process, in the AI era, algorithmic systems increasingly perform functions of analysis, prediction, and knowledge synthesis. To understand the contribution of each literature to the formation of these themes, a mapping of the study's focus, main findings, and their relevance to this research was conducted. The results of this mapping are presented in Table 1.

Table 1 literature analysis matrix

Nu.	Author (year)	Study focus	Main findings	Relevance
1	Floridi & Chiriatti (2020)	The nature of GPT and generative AI	AI generates text without full rational understanding	Epistemological basis of AI
2	Rai (2020)	Explainable AI	Transparency is a prerequisite for scientific trust	Validity of the scientific method
3	Topol (2019)	AI in healthcare	AI is effective when combined with human expertise	Human-machine integration
4	Eschenbach (2021)	Black box problem	Lack of explanation reduces trust	Epistemic problems
5	Mittelstadt (2019)	AI Ethics	Ethical principles need real implementation	Axiological dimension
6	Bender et al. (2021)	The risks of large language models	Bias and reproduction of social inequality	Criticism of data ontology
7	Dwivedi et al. (2021)	Multidisciplinary perspective of AI	AI is transforming research across fields	Transformation of science
8	Brynjolfsson et al. (2023)	Generative AI and productivity	AI improves the efficiency of scientific work	Methodological effectiveness
10	Cascella et al. (2023)	ChatGPT and medical research implications	AI opens up new opportunities but requires scientific and ethical validation	The future of AI-based research

Source: secondary data, processed

As seen in Table 1, the analyzed literature demonstrates a common trend: the development of AI has epistemological, ontological, and axiological consequences that cannot be explained solely through technical approaches. Floridi & Chiriatti (2020); Eschenbach (2021) highlight epistemological issues related to AI's ability to produce useful output without adequate rational understanding. Meanwhile, Rai (2020) emphasizes the importance of explainability as a prerequisite for scientific legitimacy. On the other hand, Mittelstadt (2019); Bender et al. (2021) note that ethical issues, data bias, and social inequality are key challenges in the implementation of AI. These findings demonstrate that the transformation of scientific methods is not only related to increasing research efficiency but also involves fundamental changes in how science is understood and practiced. To identify the most dominant thematic tendencies in the literature, the frequency of these themes across all analyzed articles was tabulated. The results are presented in Table 2.

Table 2 frequency of themes found from 10 articles

Findings theme	Number of articles	Percentage
Algorithmic transparency	9	90%
Data bias and fairness	8	80%
Human-AI integration	7	70%
Reconstruction of the scientific method	6	60%
Efficiency and productivity	6	60%
Ethics and social responsibility	8	80%

Source: secondary data, processed

Based on Table 2, algorithm transparency and explainability were the most dominant themes, appearing 90% of the time. This finding indicates that researchers' primary focus is no longer solely on improving model accuracy, but rather on AI systems' ability to explain their decision-making processes. Furthermore, issues of data bias, fairness, ethics, and social responsibility each appeared in 80% of the articles analyzed. The high frequency of these two themes indicates that AI development is increasingly understood as a social and philosophical issue, rather than simply a technical one. Meanwhile, the theme of human-AI integration appeared in 70% of the articles, indicating a growing consensus that the future of scientific research depends on collaboration between human and machine intelligence.

The first theme that emerged from the analysis was the ontological shift in the object of scientific knowledge. The results show that AI has expanded the scope of scientific ontology from directly observable empirical phenomena to a digital reality mediated by data. Of the ten articles analyzed, seven (70%) asserted that digital data, metadata, online activity traces, and computational representations are now the primary sources of scientific knowledge formation. This finding suggests that science is no longer dealing solely with tangible empirical reality, but also with a reality that has undergone a datafication process (datafied reality). Consequently, the philosophy of science needs to reexamine the ontological status of data: whether data is an objective representation of reality or a construct influenced by system design, platforms, and specific interests. Thus, AI not only expands the scope of scientific study but also changes how reality is understood in the research process.

The second theme relates to epistemological transformations in the process of knowledge formation and validation. The results show that 9 articles (90%) highlight the dominance of AI-based predictive models that produce accurate outputs without explicitly explaining causal relationships. This situation has led to a shift from an

explanatory paradigm to a predictive paradigm. In the classical scientific method, the validity of knowledge is determined by the ability to explain causal relationships that can be tested and replicated. However, in many AI applications, predictive ability often outperforms explanatory ability. This finding suggests that the new paradigm of the scientific method cannot simply rely on predictive accuracy but must also consider transparency, reproducibility, interpretability, and explainability as prerequisites for scientific legitimacy. Therefore, AI is fostering the emergence of a new form of epistemology that integrates algorithmic predictive ability with the principles of scientific rationality and openness.

The third theme shows that the axiological dimension is one of the most dominant issues in AI-based research. Eight articles (80%) emphasized that the values of utility, justice, and social responsibility should guide the use of AI. Issues such as algorithmic bias, data discrimination, privacy violations, information manipulation, and unequal access to technology were central concerns in the analyzed literature. These findings demonstrate that science is never entirely value-free. AI systems are constructed through human decisions regarding data selection, optimization objectives, model parameters, and implementation contexts. Therefore, the technical success of a system does not necessarily indicate its moral worthiness. In this context, the philosophy of science serves as a reflective framework to ensure that technological innovations are not only empirically valid but also just and beneficial to society.

The final theme suggests that the new paradigm of the scientific method does not lead to AI domination over humans, but rather to the integration of both in a collaborative model. Seven articles (70%) argue that AI is most effective when positioned as a decision-support system or cognitive partner for researchers. AI excels in processing large amounts of data, identifying hidden patterns, conducting complex simulations, and generating rapid predictions. Conversely, humans retain a strategic role in formulating research problems, determining the relevance of findings, conducting conceptual interpretations, and considering the ethical implications of research results. These findings suggest that the future of the scientific method lies in the synthesis of AI's computational capacity and human reflective wisdom. Based on these findings, this study formulates a conceptual model of a new AI-based scientific method that compares the characteristics of the classical scientific method with those of the paradigms emerging in the era of artificial intelligence. This conceptual synthesis is presented in Table 3.

Table 3 new paradigm models of AI-based scientific methods

Components	Classical scientific method	New AI-based paradigm
Object of study	Direct empirical phenomena	Empirical phenomena + digital data
Data collection	Manual/direct observation	Sensors, big data, automation
Analysis	Conventional statistics	Machine learning & AI
Validation	Hypothesis testing	Hypothesis testing + explainability
Researcher's role	Full dominance	Collaborating with AI
Core values	Objectivity	Objectivity + ethics + fairness

Source: secondary data, processed

As shown in Table 3, the transformation of scientific methods has occurred in almost all major components of research. The most significant changes are the expansion of the object of study from empirical phenomena to digital data, the use of automated systems for data collection, the application of machine learning to the analysis process, and the addition of explainability as part of the scientific validation

mechanism. Furthermore, the role of researchers has shifted from sole actors to collaborative partners with AI systems. Equally important, this new paradigm broadens the value orientation of science from mere objectivity to an integration of objectivity, ethics, fairness, and social accountability. Thus, AI does not replace the classical scientific method, but rather transforms it into a model better adapted to the complexities of the digital world while remaining rooted in the principles of the philosophy of science. Overall, the research results indicate that the transformation of scientific methods in the AI era is occurring at three levels simultaneously: ontological, epistemological, and axiological. These changes are leading to the emergence of a new scientific paradigm that positions AI as a cognitive partner in knowledge production. At the same time, the philosophy of science serves as a reflective foundation that maintains transparency, accountability, and a humanitarian orientation in contemporary scientific development.

Ontological shift: data as a new reality in science

Research findings indicate that the development of Artificial Intelligence (AI) has triggered a significant ontological transformation in science. This change is marked by a shift in the object of scientific study from directly observed empirical phenomena to a reality mediated by digital data. In the classical scientific tradition, the object of science is understood as natural, social, and human behavioral phenomena that can be observed through sensory experience or research instruments. However, in the AI era, digital data, metadata, digital footprints, sensor recordings, and computational representations have become the primary sources of knowledge (Partarakis & Zabulis, 2024). These findings indicate that the reality that forms the basis of research is no longer limited to directly observable phenomena, but also includes data generated, stored, and processed by digital systems.

These developments are inextricably linked to the increasingly comprehensive digitalization of human life. Economic activities, education, health, communication, and even social interactions now generate enormous volumes of data continuously. In this context, AI plays a central role, thanks to its ability to process data on a scale far beyond human capacity, discover patterns invisible through conventional observation, and generate predictions based on complex statistical relationships (Wang et al., 2023). Consequently, data is no longer positioned as a complement to empirical observation but rather as the primary medium for understanding and explaining various phenomena. This shift marks the birth of what can be called datafied reality, a reality understood through the process of transforming phenomena into data that can be calculated, analyzed, and modeled algorithmically.

From a philosophical perspective, this shift has profound ontological implications. The question of "what constitutes the object of science" can no longer be answered solely by referring to the empirical phenomena directly present to researchers. In the AI era, research objects often exist in digital representations that have undergone selection, categorization, and technical transformation (Kotlarsky & Oshri, 2023; Partarakis & Zabulis, 2024). Therefore, data cannot be understood as a completely neutral reflection of reality. Data are constructed and influenced by system design, information-collection mechanisms, methodological choices, and underlying institutional interests. In other words, the relationship between reality and data is mediative, not identical. Awareness of the constructive nature of data is crucial to prevent researchers from falling into the trap of assuming that data automatically represents reality objectively and completely.

The findings of this study align with those of Dwivedi et al. (2021), who asserted that AI has evolved into a multidisciplinary phenomenon transforming how organizations, governments, and the scientific community understand social reality. The advent of AI has enabled decision-making and knowledge production to rely increasingly on the analysis of large amounts of data (Zong & Guan, 2024). Thus, social reality is no longer understood solely through direct observation but also through digital representations generated by various information systems. This perspective reinforces the argument that data has acquired a new ontological status in contemporary scientific practice.

The results of this study also align with those of Bender et al. (2021), which showed that data used in AI models always carry the social, cultural, and historical biases of the environments in which they were generated. The similarity between this study and Bender et al.'s (2021) study lies in the view that data is not a completely objective and value-free entity. However, this study goes further by positioning data not merely as a potential source of bias but as a new form of reality that alters the ontological foundations of science. While previous research has focused more on dataset quality and the consequences of bias on AI model performance, this study highlights how the dominance of data in the research process shifts how science defines its objects of study.

Furthermore, Floridi & Chiriatti (2020) explain that generative AI systems generate output based on statistical pattern recognition derived from training data, rather than on an ontological understanding of the real world. This view reinforces the findings of this study that the reality processed by AI is essentially a representation of data, not reality itself. AI does not understand the meaning of objects as humans do; rather, it processes the statistical relationships contained in the data (Friedrich et al., 2022). This situation suggests that the greater AI's role in research, the more important philosophical reflection on the relationship between data, representation, and reality becomes. Thus, the primary contribution of these findings lies in affirming that AI not only brings methodological innovations in data processing but also shifts the ontological foundations of science. These changes require a reconstruction of understanding regarding the nature of research objects, the epistemic status of data, and the relationship between digital representations and empirical reality. In this context, the philosophy of science plays a crucial role in ensuring that AI developments not only enhance the analytical capabilities of research but also maintain a critical awareness of the ontological assumptions underlying the production of scientific knowledge. Therefore, the scientific paradigm in the AI era needs to view data not merely as a technical instrument but as a new ontological entity that shapes how humans understand and explain the world.

Epistemological transformation: from causal explanation to algorithmic prediction

Research results show that the development of Artificial Intelligence (AI) has driven a fundamental epistemological transformation in the practice of science. A shift in the source of knowledge legitimacy from the dominance of causal explanations to the increasing role of algorithmic predictions marks this change. In the tradition of the classical scientific method, knowledge is considered valid if it can explain causal relationships logically and empirically and can be tested through processes of verification and falsification (Yao, 2024). Scientific explanations are required not only to provide answers about what happens, but also why a phenomenon occurs. Thus, causal explanations are the primary foundation in the formation of scientific knowledge.

However, the results of this study indicate that this paradigm is beginning to shift with the increasing use of AI in various research fields. Machine learning and deep learning-based systems can produce predictions with very high accuracy, even though the mechanisms that generate them are often difficult to explain explicitly (Taye, 2023). Unlike conventional scientific approaches that rely on theory building and hypothesis testing, AI algorithms identify statistical patterns, complex correlations, and nonlinear relationships in large-scale datasets. In many cases, AI models can predict disease risks, consumer behavior, market changes, and even environmental phenomena with a level of accuracy that surpasses traditional analytical approaches, even without providing adequate causal explanations for the factors underlying these predictions.

Epistemologically, this situation indicates the emergence of a new form of knowledge oriented toward predictive knowledge. Knowledge is no longer solely assessed by its ability to explain reality, but also by its ability to anticipate and predict future events (Hussien et al., 2025). This shift can be understood as a response to the complexity of contemporary phenomena, which are increasingly difficult to reduce to simple cause-and-effect relationships. Various modern issues, such as the spread of epidemics, global market dynamics, climate change, user behavior in digital media, and interconnected social systems, generate enormous volumes of data and complex interaction patterns. In such situations, AI's predictive capabilities are often more effective than conventional causal approaches in generating relevant information for decision-making.

Nevertheless, the findings of this study indicate that predictive accuracy cannot be the sole basis for scientific legitimacy. Knowledge generated by AI must still meet the basic principles of scientific epistemology, namely transparency of the inference process, reproducibility of results, interpretability, and explainability of the rationale behind a decision (Hussien et al., 2025). Without these elements, AI predictions risk losing legitimacy because the scientific community cannot critically evaluate them. Therefore, the epistemological transformation that occurs is not a total replacement of explanation by prediction, but rather an expansion of the epistemological framework that accommodates both as complementary sources of knowledge. In this new paradigm, prediction enables understanding of future possibilities, while explanation remains necessary to provide rational justification for the results obtained.

The findings of this study align with Rai (2020), who asserted that explainable AI is a crucial prerequisite for building trust in algorithm-based systems. According to Rai, the ability to explain the decision-making process is a key factor in determining whether AI results are acceptable in contexts that demand high accountability. This perspective reinforces the findings of this study, which argue that the legitimacy of scientific knowledge is inextricably linked to the system's ability to explain the rationale behind a conclusion. A similar argument is also put forward by von Eschenbach (2021) in his discussion of the black box problem. He asserts that the primary problem with AI lies not solely in the possibility of technical errors, but in humans' inability to understand the inference process that produces algorithmic decisions. This research, like this study, holds that high accuracy does not automatically lead to scientific trust. Instead, valid knowledge still requires justification that can be understood and rationally evaluated. Thus, transparency is an inseparable epistemological component of AI use in scientific research.

In a more practical context, Topol's (2019) research in the healthcare sector showed that physicians' acceptance of AI systems increased when their recommendations could be explained and verified through clinical reasoning. These

findings indicate that the relationship between humans and AI is not solely determined by technological performance, but also by the system's ability to interact with existing human knowledge frameworks. The results of this study reinforce this view by demonstrating that explainability serves as a bridge between AI's computational efficiency and prevailing standards of scientific validity within the academic community.

While sharing some similarities with previous research, this study offers a different perspective. Most previous research positions transparency and explainability as practical necessities for increasing user trust in AI across sectors such as healthcare, business, and government. In contrast, this study positions these issues within a more fundamental philosophical framework. Transparency is understood not merely as a technical feature or instrument of technology governance, but as an epistemological requirement that determines whether a result can be recognized as valid scientific knowledge. Thus, the focus of this research is not limited to how AI is used, but rather on how AI changes fundamental concepts of validity, objectivity, and justification of knowledge.

The primary contribution of this finding lies in its attempt to reconcile two often-conflicting epistemological traditions: the rational-explanation tradition and the computational-prediction tradition. The results suggest that the future of the scientific method lies not in the dominance of one approach, but in the integration of both within an epistemological framework better adapted to technological developments. Within this paradigm, AI's predictive capabilities enhance efficiency and accuracy in understanding complex phenomena. At the same time, the principles of scientific explanation continue to serve as a justification mechanism, ensuring the rationality, transparency, and accountability of knowledge. Therefore, the epistemological transformation in the AI era does not mark the end of the classical scientific method, but rather its evolution toward a form more responsive to the challenges of contemporary science.

Axiological challenges: ethics, justice, and social responsibility

The research findings show that the transformation of scientific methods in the era of Artificial Intelligence (AI) not only brings about ontological and epistemological changes but also presents increasingly complex axiological challenges. The axiological dimension is an inseparable aspect of the use of AI in research because this technology not only influences how knowledge is produced but also determines how that knowledge is used and impacts society. The research findings show that the use of AI in science offers various benefits, such as increased analytical efficiency, the ability to process large-scale data, accelerated decision-making, and expanded research capacity. However, behind these benefits lie risks such as algorithmic bias, decision-making discrimination, privacy violations, a lack of accountability, and unequal access to technology. This situation shows that the technical success of AI is not always directly proportional to its moral and social acceptability.

Philosophically, these findings indicate that, in the AI era, science is increasingly difficult to view as value-free. AI systems are built through a series of human decisions spanning data selection, optimization objectives, algorithm design, training parameters, and implementation context (Gupta et al., 2022). Each of these decisions contains specific assumptions, preferences, and values that potentially influence the system's outcomes. Thus, AI is never truly neutral, reflecting the characteristics of the data and the choices made by its developers and the institutions that use it. In many cases, the datasets used to train AI models contain historical biases stemming from unequal social

structures. As a result, AI systems can reproduce and even reinforce pre-existing inequities if not designed and supervised critically.

From a philosophical perspective, this issue is directly related to the axiological dimension, which questions the purpose, benefits, and consequences of using scientific knowledge. Axiology not only assesses whether knowledge is empirically true, but also whether it is used for good purposes and provides equitable benefits to society (Setiawan et al., 2025). In the context of AI, axiological questions become increasingly important because decisions made by algorithmic systems can directly impact human lives, from health diagnoses and job selection to credit assessments and education systems, to public policy (Gomez, 2025). Therefore, evaluation of AI cannot be based solely on accuracy or efficiency; it must also consider its impact on justice, social welfare, individual rights, and human dignity.

The rationale for these findings lies in the fact that science always operates in a social space fraught with vested interests and practical consequences. The greater the influence of AI on decision-making processes, the greater the need to ensure that the technology is used responsibly (Khalid et al., 2024). Within this framework, research shows that ethical considerations cannot be placed as a final step after the technology has been developed. Instead, ethics needs to be integrated from the formulation of the research problem, data collection, model design, interpretation of results, and even technology implementation. This approach positions ethics as an inherent part of the scientific method, not simply a corrective tool when problems arise. Thus, a highly accurate AI system that discriminates against or harms certain groups cannot be viewed as a fully successful form of scientific progress.

The findings of this study align with Mittelstadt's (2019) view, which asserts that AI ethical principles will be ineffective without concrete implementation mechanisms. According to Mittelstadt (2019), principles such as fairness, transparency, accountability, and nonmaleficence often remain at the normative level without being translated into operational procedures that can be implemented in practice. This perspective reinforces the findings of this study, which argue that ethics must be present not only as a declaration of values but also as part of the scientific process and technology governance. In other words, the success of AI is measured not only by what a system can do, but also by how responsibly it is designed and used.

The results of this study are also consistent with those of Bender et al. (2021), who found that large-scale language models have the potential to reproduce stereotypes, prejudices, and social inequalities contained in their training data. The similarity between this study and Bender et al.'s (2021) study lies in their rejection of the assumption that technology is a completely neutral instrument. Both studies emphasize that AI systems always carry the social and historical imprint of the data they use. However, this study broadens this discussion by positioning the issue of bias not merely as a technical problem to be addressed through improved data quality, but as an axiological issue related to moral responsibility in the production of scientific knowledge.

Furthermore, research by Cascella et al. (2023) in the health sector shows that AI has great potential to improve the quality of diagnosis, clinical decision-making, and the effectiveness of healthcare services. However, implementing this technology still requires rigorous scientific validation and a robust ethical framework, as it concerns human safety and well-being. These findings reinforce this study's conclusion that the greater an AI system's impact on human life, the greater the demand for its ethical accountability. Therefore, the success of AI implementation depends not only on the

technology's performance but also on the system's ability to meet applicable moral and social standards.

While related to various previous studies, this study offers a different perspective on the relationship between ethics and the scientific method. Most previous research positions ethics as an external regulatory instrument that controls the negative impacts of technology after it has been developed. In contrast, this study views ethics as an internal component inherent in the knowledge production process itself. From this perspective, questions of fairness, responsibility, and social utility must be part of the scientific process from the outset, as important as questions of data validity or research accuracy. The primary contribution of these findings lies in broadening the understanding of scientific objectivity in the AI era. Objectivity can no longer be defined narrowly as the ability to produce accurate knowledge free from empirical error. Instead, objectivity needs to be understood more comprehensively, incorporating elements of social responsibility, distributive justice, protection of individual rights, and respect for human values. Thus, the new paradigm of AI-based scientific methods demands not only empirical truth and technological efficiency but also a strong ethical commitment to ensure that scientific development remains oriented toward human well-being and social sustainability. In this context, the philosophy of science serves as a reflective foundation, ensuring that technological innovation does not lose its normative direction amid the acceleration of digital transformation.

Human-machine integration as a new paradigm of scientific method

The research findings indicate that the development of Artificial Intelligence (AI) does not lead to the replacement of humans in the scientific process, but rather to a pattern of increasingly close collaboration between human and machine intelligence. In this paradigm, AI functions as a system that excels at large-scale data processing, complex pattern identification, simulation of phenomena, and accelerated prediction. Meanwhile, humans continue to play an irreplaceable role in conceptual and normative aspects, such as formulating research problems, interpreting analytical results, validating the meaning of findings, and determining the direction and ethical implications of the research. These findings demonstrate that the division of roles between humans and machines is not competitive, but rather complementary.

Rationally, this situation arises from fundamental differences between AI's computational capacity and human cognitive capacity. AI excels in speed, scale, and consistency in processing large amounts of data, but lacks the reflective consciousness, moral intuition, or contextual understanding inherent in human experience (Kim & Lee, 2025). In contrast, humans possess the ability to make normative judgments, understand social context, and provide meaningful interpretations of analytical results. However, they are limited in their ability to handle the complexity of vast, multidimensional data. Therefore, integrating the two produces a form of scientific method better suited to the complexity of contemporary problems. Within this framework, AI is not positioned as a replacement for scientists, but rather as a cognitive partner that expands thinking capacity and strengthens scientific decision-making processes.

These findings align with Topol's (2019) observation that, in a medical context, AI performance is optimal when used as a clinical decision-support tool rather than as a substitute for healthcare professionals. In this context, AI accelerates medical data analysis, while the final decision rests with healthcare professionals, who consider both clinical and ethical aspects. Thus, collaboration between humans and machines has been

shown to improve outcome quality without diminishing the role of human judgment. This observation reinforces research findings that the primary value of AI lies in its ability to work synergistically with humans in a shared decision-making system.

A similar view was expressed by Brynjolfsson et al. (2023), who found that the use of generative AI can significantly increase work productivity, especially when human users continue to monitor, evaluate, and validate the system's results. These findings suggest that the greatest benefits of AI do not come from full automation, but rather from the controlled integration of algorithmic capabilities and human judgment. In this context, AI acts as an augmentation tool that extends human capabilities, rather than completely replacing them. This aligns with the findings of this study, which places human-machine collaboration at the heart of the transformation of the scientific method in the digital age.

Furthermore, a study by Jeblick et al. (2023) showed that using AI, such as ChatGPT, in a radiology context can simplify medical reports and improve the efficiency of scientific communication. However, the study also emphasized that the use of AI still requires human verification to ensure accuracy, information security, and appropriateness to the clinical context. Similarities with this study lie in recognizing that AI significantly improves the efficiency and quality of scientific work but cannot completely replace the role of humans in validation and final decision-making.

While sharing similarities with previous research, this study offers a broader perspective on understanding human-machine integration. Most previous research tends to view human-AI collaboration instrumentally, as a means to improve efficiency, productivity, or accuracy in a specific field. In contrast, this study views this integration as a paradigmatic transformation within the scientific method itself. This means that human-machine collaboration not only impacts technical performance but also changes the fundamental structure of scientific practice, including how knowledge is generated, validated, and interpreted. The primary contribution of these findings lies in affirming that the new paradigm of AI-based scientific methods is not simply the result of technological adoption, but rather a form of epistemic reorganization between humans and machines in the production of knowledge. In this paradigm, science is no longer understood as the exclusive product of human labor, but rather as the product of a dynamic interaction between AI's computational capacity and human reflective capacity. Therefore, the philosophy of science continues to play a crucial role as a normative framework that ensures this integration occurs within the boundaries of rationality, accountability, and ethical orientation. Thus, human-machine collaboration should not only enhance scientific efficiency but also remain directed toward strengthening the fundamental goals of science: a better understanding of reality and the enhancement of human well-being.

Conclusions

The development of Artificial Intelligence (AI) has brought about a fundamental transformation in the scientific method, encompassing three main dimensions of the philosophy of science: ontological, epistemological, and axiological, and giving rise to the paradigm of human-machine integration. Ontologically, the object of science shifts from direct empirical phenomena to a digital data-based reality mediated by computational systems. Epistemologically, there has been a shift from the dominance of causal explanations to the strengthening of algorithmic predictions based on complex correlations, although still requiring the principles of scientific transparency and explainability. Axiologically, the use of AI presents significant ethical challenges, such as

bias, injustice, and social inequality, which demand integrating values into the entire scientific process. On the other hand, the scientific method in the AI era does not lead to human substitution, but rather to a synergistic collaboration between human and machine intelligence in the production of scientific knowledge.

The implications of these findings emphasize that AI plays a role not only as a technical tool but also as an epistemic force transforming the fundamental structure of modern science. The main contribution of this research lies in developing a conceptual framework that integrates ontology, epistemology, and axiology to interpret the transformation of scientific methods in the AI era, while also emphasizing the importance of the philosophy of science as a normative foundation for maintaining rationality, transparency, and social responsibility in AI-based research. Thus, scientific objectivity is no longer understood narrowly as value-neutrality, but rather as a balance among accuracy, explainability, and social justice in knowledge production.

This research has limitations because it is entirely based on the literature and draws on a limited number of sources, so it does not directly reflect empirical dynamics across various contexts of scientific practice. Furthermore, the analysis is conceptual and has not been tested through empirical approaches or case studies in specific fields. Therefore, further research is recommended using a mixed-methods approach or empirical case studies in specific fields such as health, education, or public policy to test the validity of the conceptual framework. Future research could also further explore how the implementation of AI in scientific practice affects real-world decision-making processes and how human-AI collaboration models can be optimized in various disciplines.

References

- Alvarado, R. (2023). AI as an Epistemic Technology. *Science and Engineering Ethics*, 29(5), 32. <https://doi.org/10.1007/s11948-023-00451-3>
- Aspers, P., & Corte, U. (2021). What is Qualitative in Research. *Qualitative Sociology*, 44(4), 599–608. <https://doi.org/10.1007/s11133-021-09497-w>
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the Dangers of Stochastic Parrots. *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. <https://doi.org/10.1145/3442188.3445922>
- Braun, V., & Clarke, V. (2021). To saturate or not to saturate? Questioning data saturation as a useful concept for thematic analysis and sample-size rationales. *Qualitative Research in Sport, Exercise and Health*, 13(2), 201–216. <https://doi.org/10.1080/2159676X.2019.1704846>
- Brynjolfsson, E., Li, D., & Raymond, L. (2023). *Generative AI at Work*. <https://doi.org/10.3386/w31161>
- Campbell, S., Greenwood, M., Prior, S., Shearer, T., Walkem, K., Young, S., Bywaters, D., & Walker, K. (2020). Purposive sampling: complex or simple? Research case examples. *Journal of Research in Nursing*, 25(8), 652–661. <https://doi.org/10.1177/1744987120927206>
- Cascella, M., Montomoli, J., Bellini, V., & Bignami, E. (2023). Evaluating the Feasibility of ChatGPT in Healthcare: An Analysis of Multiple Clinical and Research Scenarios. *Journal of Medical Systems*, 47(1), 33. <https://doi.org/10.1007/s10916-023-01925-4>
- Dwivedi, Y. K., Hughes, L., Ismagilova, E., Aarts, G., Coombs, C., Crick, T., Duan, Y., Dwivedi, R., Edwards, J., Eirug, A., Galanos, V., Ilavarasan, P. V., Janssen, M., Jones, P., Kar, A.

- K., Kizgin, H., Kronemann, B., Lal, B., Lucini, B., ... Williams, M. D. (2021). Artificial Intelligence (AI): Multidisciplinary perspectives on emerging challenges, opportunities, and agenda for research, practice and policy. *International Journal of Information Management*, 57, 101994. <https://doi.org/10.1016/j.ijinfomgt.2019.08.002>
- Eschenbach, W. J. von. (2021). Transparency and the Black Box Problem: Why We Do Not Trust AI. *Philosophy & Technology*, 34(4), 1607–1622. <https://doi.org/10.1007/s13347-021-00477-0>
- Floridi, L., & Chiriatti, M. (2020). GPT-3: Its Nature, Scope, Limits, and Consequences. *Minds and Machines*, 30(4), 681–694. <https://doi.org/10.1007/s11023-020-09548-1>
- Friedrich, S., Antes, G., Behr, S., Binder, H., Brannath, W., Dumpert, F., Ickstadt, K., Kestler, H. A., Lederer, J., Leitgöb, H., Pauly, M., Steland, A., Wilhelm, A., & Friede, T. (2022). Is there a role for statistics in artificial intelligence? *Advances in Data Analysis and Classification*, 16(4), 823–846. <https://doi.org/10.1007/s11634-021-00455-6>
- Gomez, E. A. R. (2025). Systematic and axiological capacities in artificial intelligence applied in the public sector. *Public Policy and Administration*, 40(2), 351–371. <https://doi.org/10.1177/09520767231170321>
- Gupta, S., Modgil, S., Bhattacharyya, S., & Bose, I. (2022). Artificial intelligence for decision support systems in the field of operations research: review and future scope of research. *Annals of Operations Research*, 308(1–2), 215–274. <https://doi.org/10.1007/s10479-020-03856-6>
- Hassija, V., Chamola, V., Mahapatra, A., Singal, A., Goel, D., Huang, K., Scardapane, S., Spinelli, I., Mahmud, M., & Hussain, A. (2024). Interpreting Black-Box Models: A Review on Explainable Artificial Intelligence. *Cognitive Computation*, 16(1), 45–74. <https://doi.org/10.1007/s12559-023-10179-8>
- Hussien, M. M., Melo, A. N., Ballardini, A. L., Maldonado, C. S., Izquierdo, R., & Sotelo, M. Á. (2025). RAG-based explainable prediction of road users behaviors for automated driving using knowledge graphs and large language models. *Expert Systems with Applications*, 265, 125914. <https://doi.org/10.1016/j.eswa.2024.125914>
- Jeblick, K., Schachtner, B., Dextl, J., Mittermeier, A., Stüber, A. T., Topalis, J., Weber, T., Wesp, P., Sabel, B. O., Ricke, J., & Ingrisich, M. (2023). ChatGPT makes medicine easy to swallow: an exploratory case study on simplified radiology reports. *European Radiology*, 34(5), 2817–2825. <https://doi.org/10.1007/s00330-023-10213-1>
- Khalid, J., Chuanmin, M., Altaf, F., Shafqat, M. M., Khan, S. K., & Ashraf, M. U. (2024). AI-Driven Risk Management and Sustainable Decision-Making: Role of Perceived Environmental Responsibility. *Sustainability*, 16(16), 6799. <https://doi.org/10.3390/su16166799>
- Kim, H., & Lee, H. sun. (2025). The Role of Artificial Intelligence in Backing Up Human Insight: A Powerful Partnership. *MechEcology*, 4(1), 1–17. <https://doi.org/10.23104/ME.e7>
- Kotlarsky, J., & Oshri, I. (2023). A Paradigm Shift in Understanding Digital Objects in IS: A Semiotic Perspective on Artificial Intelligence Technologies. In *Advancing Information Systems Theories, Volume II* (pp. 119–148). https://doi.org/10.1007/978-3-031-38719-7_4
- Kumar, A., Dhanka, S., Sharma, A., Sharma, A., Nain, M., Kumar, P., Gupta, A. R., Bansal, J., Saxena, N. K., & Pant, R. (2026). A Comprehensive Review of Bias in AI, ML, and DL Models: Methods, Impacts, and Future Directions. *Archives of Computational*

- Methods in Engineering*, 33(4), 5743–5773. <https://doi.org/10.1007/s11831-025-10483-6>
- Mittelstadt, B. (2019). Principles alone cannot guarantee ethical AI. *Nature Machine Intelligence*, 1(11), 501–507. <https://doi.org/10.1038/s42256-019-0114-4>
- Nowell, L. S., Norris, J. M., White, D. E., & Moules, N. J. (2017). Thematic Analysis. *International Journal of Qualitative Methods*, 16(1). <https://doi.org/10.1177/1609406917733847>
- Partarakis, N., & Zabulis, X. (2024). A Review of Immersive Technologies, Knowledge Representation, and AI for Human-Centered Digital Experiences. *Electronics*, 13(2), 269. <https://doi.org/10.3390/electronics13020269>
- Rai, A. (2020). Explainable AI: from black box to glass box. *Journal of the Academy of Marketing Science*, 48(1), 137–141. <https://doi.org/10.1007/s11747-019-00710-5>
- Razzaq, K., & Shah, M. (2025). Machine Learning and Deep Learning Paradigms: From Techniques to Practical Applications and Research Frontiers. *Computers*, 14(3), 93. <https://doi.org/10.3390/computers14030093>
- Ren, M., Chen, N., & Qiu, H. (2023). Human-machine Collaborative Decision-making: An Evolutionary Roadmap Based on Cognitive Intelligence. *International Journal of Social Robotics*, 15(7), 1101–1114. <https://doi.org/10.1007/s12369-023-01020-1>
- Setiawan, A., Syukri SS, A., & Yenti, Z. (2025). Axiology of Science and Its Benefits for Humanity. *Arus Jurnal Sosial Dan Humaniora*, 5(2), 2247–2254. <https://doi.org/10.57250/ajsh.v5i2.1418>
- Shah, M., & Sureja, N. (2025). A Comprehensive Review of Bias in Deep Learning Models: Methods, Impacts, and Future Directions. *Archives of Computational Methods in Engineering*, 32(1), 255–267. <https://doi.org/10.1007/s11831-024-10134-2>
- Sreerama, J., & Krishnamoorthy, G. (2022). Ethical Considerations in AI Addressing Bias and Fairness in Machine Learning Models. *Journal of Knowledge Learning and Science Technology ISSN: 2959-6386 (Online)*, 1(1), 130–138. <https://doi.org/10.60087/jklst.vol1.n1.p138>
- Taye, M. M. (2023). Understanding of Machine Learning with Deep Learning: Architectures, Workflow, Applications and Future Directions. *Computers*, 12(5), 91. <https://doi.org/10.3390/computers12050091>
- Topol, E. J. (2019). High-performance medicine: the convergence of human and artificial intelligence. *Nature Medicine*, 25(1), 44–56. <https://doi.org/10.1038/s41591-018-0300-7>
- Tsvetkova, M., Yasseri, T., Pescetelli, N., & Werner, T. (2024). A new sociology of humans and machines. *Nature Human Behaviour*, 8(10), 1864–1876. <https://doi.org/10.1038/s41562-024-02001-8>
- Wang, H., Fu, T., Du, Y., Gao, W., Huang, K., Liu, Z., Chandak, P., Liu, S., Van Katwyk, P., Deac, A., Anandkumar, A., Bergen, K., Gomes, C. P., Ho, S., Kohli, P., Lasenby, J., Leskovec, J., Liu, T.-Y., Manrai, A., ... Zitnik, M. (2023). Scientific discovery in the age of artificial intelligence. *Nature*, 620(7972), 47–60. <https://doi.org/10.1038/s41586-023-06221-2>
- Xiao, Y., & Watson, M. (2019). Guidance on Conducting a Systematic Literature Review. *Journal of Planning Education and Research*, 39(1), 93–112. <https://doi.org/10.1177/0739456X17723971>
- Yao, Q. (2024). Concepts and Reasoning: a Conceptual Review and Analysis of Logical Issues in Empirical Social Science Research. *Integrative Psychological and*

- Behavioral Science*, 58(2), 502–530. <https://doi.org/10.1007/s12124-023-09792-x>
- Zhu, S., Yu, T., Xu, T., Chen, H., Dustdar, S., Gigan, S., Gunduz, D., Hossain, E., Jin, Y., Lin, F., Liu, B., Wan, Z., Zhang, J., Zhao, Z., Zhu, W., Chen, Z., Durrani, T. S., Wang, H., Wu, J., ... Pan, Y. (2023). Intelligent Computing: The Latest Advances, Challenges, and Future. *Intelligent Computing*, 2. <https://doi.org/10.34133/icomputing.0006>
- Zong, Z., & Guan, Y. (2024). AI-Driven Intelligent Data Analytics and Predictive Analysis in Industry 4.0: Transforming Knowledge, Innovation, and Efficiency. *Journal of the Knowledge Economy*, 16(1), 864–903. <https://doi.org/10.1007/s13132-024-02001-z>