



Abdelkader Ouanouki

E-mail: Aekouanouki@gmail.com

Orcid: <https://orcid.org/0009-0005-1394-446X>

Djamel Ouanouki

E-mail: djamal1313@yahoo.com

Ziane Achour University of Djelfa, Algeria

Cita sugerida (APA, séptima edición)

Ouanouki, A., & Ouanouki, D. (2026). Artificial intelligence from sociological and legal perspectives. *Revista Sociedad & Tecnología*, 9(S2), 885-895, DOI: <https://doi.org/10.51247/st.v9iS2.796>

==== o ====

Artificial intelligence from sociological and legal perspectives

ABSTRACT

This study examined artificial intelligence from sociological and legal perspectives, with the objective of analyzing its principal impacts on social structures, labor relations, rights protection, and regulatory systems. A qualitative methodology with a documentary-descriptive design was adopted. Scientific articles, academic books, and legal sources indexed in recognized databases were reviewed through thematic content analysis. The results showed that artificial intelligence has transformed productive processes, decision-making systems, and digital interactions, while also generating challenges related to unemployment, inequality, privacy, algorithmic bias, and legal accountability. From the sociological perspective, AI was found to reshape identities, communication practices, and access to opportunities. From the legal perspective, existing frameworks still face difficulties in regulating autonomous systems, assigning responsibility, and protecting personal data. It was concluded that artificial intelligence should not be addressed exclusively as a technological issue, but as a multidimensional phenomenon requiring interdisciplinary governance. A human-centered regulatory model is necessary to balance innovation with social justice, transparency, and democratic values. The study emphasized that the future effects of AI will depend on the ethical and institutional decisions adopted by governments, organizations, and societies in the coming years.

Keywords: artificial intelligence, sociology, law, regulation.

==== o ====

Inteligencia artificial desde las perspectivas sociológica y jurídica

RESUMEN

El presente estudio examinó la inteligencia artificial desde las perspectivas sociológica y jurídica, con el objetivo de analizar sus principales impactos en las estructuras sociales, las relaciones laborales, la protección de derechos y los sistemas regulatorios. Se adoptó una metodología cualitativa con diseño documental-descriptivo. Se revisaron artículos científicos, libros académicos y fuentes jurídicas indexadas en bases de datos reconocidas mediante análisis temático de contenido. Los resultados evidenciaron que la inteligencia artificial ha transformado los procesos productivos, los sistemas de toma de decisiones y las interacciones

digitales, al tiempo que genera desafíos vinculados al desempleo, la desigualdad, la privacidad, el sesgo algorítmico y la responsabilidad legal. Desde la perspectiva sociológica, se determinó que la IA reconfigura identidades, prácticas comunicativas y acceso a oportunidades. Desde la perspectiva jurídica, los marcos normativos actuales aún presentan dificultades para regular sistemas autónomos, asignar responsabilidades y proteger datos personales. Se concluyó que la inteligencia artificial no debe abordarse exclusivamente como un asunto tecnológico, sino como un fenómeno multidimensional que requiere gobernanza interdisciplinaria. Se necesita un modelo regulatorio centrado en el ser humano que equilibre innovación con justicia social, transparencia y valores democráticos.

Palabras clave: inteligencia artificial, sociología, derecho, regulación.

==== o ====

Inteligência artificial sob perspectivas sociológicas e jurídicas

RESUMO

Este estudo examinou a inteligência artificial sob perspectivas sociológicas e jurídicas, com o objetivo de analisar seus principais impactos nas estruturas sociais, nas relações de trabalho, na proteção de direitos e nos sistemas regulatórios. Foi adotada uma metodologia qualitativa com desenho documental-descritivo. Artigos científicos, livros acadêmicos e fontes jurídicas indexadas em bases de dados reconhecidas foram revisados por meio de análise temática de conteúdo. Os resultados mostraram que a inteligência artificial transformou processos produtivos, sistemas de tomada de decisão e interações digitais, ao mesmo tempo em que gera desafios relacionados ao desemprego, desigualdade, privacidade, viés algorítmico e responsabilidade jurídica. Sob a perspectiva sociológica, verificou-se que a IA reconfigura identidades, práticas comunicativas e acesso a oportunidades. Sob a perspectiva jurídica, os marcos normativos atuais ainda enfrentam dificuldades para regular sistemas autônomos, atribuir responsabilidades e proteger dados pessoais. Concluiu-se que a inteligência artificial não deve ser tratada exclusivamente como questão tecnológica, mas como fenômeno multidimensional que exige governança interdisciplinar. Torna-se necessário um modelo regulatório centrado no ser humano que equilibre inovação, justiça social, transparência e valores democráticos.

Palavras-chave: inteligência artificial, sociologia, direito, regulação.

==== o ====

INTRODUCTION

Artificial intelligence (AI) has become one of the most transformative forces of the twenty-first century, influencing economic production, governance, communication, and everyday decision-making. What began as a specialized branch of computer science has evolved into a multidisciplinary phenomenon with consequences that extend far beyond technology itself. As noted by Yadav et al. (2020), AI systems are designed to perform tasks traditionally associated with human intelligence, including learning, reasoning, perception, and problem-solving. This rapid expansion has generated both enthusiasm and concern, particularly regarding its effects on institutions, labor markets, and legal systems. Therefore, understanding AI requires approaches that integrate technical innovation with social and normative analysis.

The historical development of AI demonstrates how quickly experimental concepts can become embedded in daily life. Early milestones, such as the Dartmouth Conference of 1956, established AI as a formal academic field and inspired decades of research into machine reasoning and automation. According to McCarthy et al. (2006), the original ambition was to simulate aspects of human intelligence through computational systems. Today, these ambitions are visible in search engines, recommendation systems, autonomous vehicles, and generative language models. The transition from theory to practical application has intensified

debates concerning accountability, ethics, and the future relationship between humans and intelligent machines.

From a sociological perspective, AI is not merely a neutral tool but a force capable of reshaping social structures and patterns of interaction. Castells (2010) argued that digital technologies transform the organization of societies by creating networked systems of communication and power. In this context, AI amplifies those transformations by mediating access to information, employment opportunities, and public services. Algorithms increasingly determine what citizens read, buy, or believe, influencing behavior in subtle but significant ways. Consequently, AI has become an important object of sociological inquiry because it changes how identities are formed and how inequalities are reproduced or challenged.

One of the most significant concerns relates to employment and social stratification. Brynjolfsson and McAfee (2014) observed that automation technologies can increase productivity while simultaneously displacing workers whose tasks become machine-performable. AI extends this tendency into cognitive professions, affecting sectors such as finance, law, education, and healthcare. While new occupations emerge in data science and digital engineering, many traditional roles face uncertainty or transformation. This uneven distribution of benefits may deepen existing inequalities if education systems and labor policies fail to adapt. Therefore, AI raises essential questions about fairness, redistribution, and the future of work.

Another major sociological issue is algorithmic bias and surveillance. O'Neil (2016) demonstrated that automated decision systems may reproduce discrimination when trained on biased historical data. Hiring platforms, predictive policing tools, and credit-scoring models can unintentionally reinforce racial, gender, or class inequalities. At the same time, AI-powered surveillance systems enable unprecedented monitoring of citizens through facial recognition, behavioral tracking, and predictive analytics. These developments challenge democratic freedoms and the right to privacy, making regulation increasingly urgent. Thus, technological efficiency must be balanced with safeguards for dignity and civil liberties.

From a legal perspective, AI creates complex challenges because existing legal systems were designed around human actors and conventional products. When autonomous systems make decisions that cause harm, responsibility becomes difficult to assign. According to Abbott (2020), questions of liability arise in cases involving self-driving vehicles, medical diagnostics, and automated financial systems. Should accountability fall on developers, operators, owners, or the algorithm itself? Current legal frameworks often struggle to address these scenarios, creating what scholars describe as an accountability gap. As a result, lawmakers worldwide are seeking adaptive regulatory models.

Data protection and intellectual property have also emerged as central legal concerns. The large-scale training of AI models often depends on personal information or copyrighted material collected from digital environments. Wachter et al. (2017) emphasized the growing demand for transparency, explainability, and user rights in algorithmic governance. Simultaneously, disputes over AI-generated texts, images, and inventions have raised questions about authorship and ownership. These controversies illustrate that AI challenges foundational legal concepts regarding consent, creativity, and property. Hence, modern legal systems must evolve to preserve justice while encouraging innovation.

Ultimately, artificial intelligence must be examined through both sociological and legal lenses because its impact is simultaneously social and normative. Sociological analysis explains how AI transforms behavior, institutions, and inequality, while legal analysis defines the rights, responsibilities, and limits necessary for democratic governance. The future of AI will depend not only on technical capability but also on the ethical and institutional choices societies make. A human-centered framework is therefore essential to ensure that AI promotes inclusion, accountability, and respect for fundamental rights in an increasingly automated world.

Methodology

This study adopted a qualitative research approach, as this methodology is appropriate for analyzing complex social and legal phenomena that require interpretation rather than numerical measurement. Qualitative research made it possible to examine artificial intelligence from sociological and legal perspectives through the analysis of concepts, theories, and regulatory debates. According to Espinoza Freire (2020a), qualitative inquiry facilitates a deeper understanding of realities shaped by human interaction, values, and institutional dynamics. Therefore, this approach was considered suitable for exploring how AI influences society and law in contemporary contexts.

The research design was documentary and descriptive in nature. A documentary method was employed because the study relied on scientific books, indexed journal articles, legal reports, and institutional documents related to artificial intelligence. Descriptive research was useful because it enabled the systematic presentation of the principal characteristics, challenges, and implications of AI in both fields of analysis. Espinoza Freire (2020b) stated that documentary research strengthens academic rigor when reliable sources are selected and critically interpreted. For this reason, specialized literature constituted the main empirical basis of the present investigation.

The process of information retrieval was carried out through academic databases recognized for scientific reliability, including Google Scholar, Scopus, Web of Science, and institutional repositories. Search strategies included the use of keywords such as "artificial intelligence," "algorithmic bias," "AI regulation," "digital society," and "legal liability." Boolean operators were also applied to refine results and improve relevance. As indicated by Espinoza Freire (2020c), the organized search for scientific information in academic databases is essential to ensure the quality, validity, and relevance of research evidence. Consequently, only peer-reviewed and academically recognized materials were prioritized.

The inclusion criteria considered publications in English and Spanish that addressed artificial intelligence from legal, sociological, ethical, or interdisciplinary perspectives. Priority was given to recent sources, although seminal works were also incorporated when they provided theoretical foundations. Exclusion criteria eliminated duplicated records, non-academic webpages, opinion pieces without scholarly support, and sources lacking methodological transparency. According to Espinoza Freire and Calva Nagua (2020), ethical rigor in research also requires transparency in source selection and responsible management of information. Thus, the review process sought coherence, traceability, and academic integrity.

For data analysis, a thematic content analysis technique was applied. Selected documents were read critically and organized into analytical categories such as labor transformation, digital inequality, privacy, algorithmic discrimination, legal accountability, and governance models. This categorization allowed the identification of recurring patterns and conceptual relationships among authors. Espinoza Freire and Rad Camayd (2020) emphasized that systematic qualitative analysis helps interpret meanings and establish coherent explanatory frameworks. Accordingly, the findings were structured around the most relevant socio-legal dimensions of AI.

Scientific argumentation guided the interpretation and discussion of results. Contrasting viewpoints were compared in order to identify consensus, controversies, and unresolved debates in the literature. This procedure strengthened analytical depth and supported evidence-based conclusions. Espinoza-Freire (2021) noted that argumentation is a fundamental didactic and scientific tool because it enables researchers to justify positions through logical reasoning and empirical support. Therefore, the discussion section was developed through comparative reflection rather than simple descriptive summary.

Ethical principles were observed throughout the study. Since the research did not involve human participants or confidential data, no experimental risk was generated. Nevertheless, intellectual honesty was maintained through proper citation practices, faithful interpretation

of authors' ideas, and avoidance of plagiarism. Espinoza-Freire (2022) affirmed that ethics in scientific research requires responsibility, transparency, and respect for knowledge production. In the same way, theoretical rigor and methodological coherence were preserved during all stages of the investigation.

Finally, the study was supported by an interpretive framework consistent with grounded theoretical reasoning, where concepts emerged from the systematic comparison of sources and categories. Although no field interviews were conducted, the progressive coding of literature allowed the integration of diverse perspectives into a coherent explanatory model. Espinoza-Freire (2024) explained that grounded analytical logic contributes to theory building through constant comparison and conceptual refinement. Consequently, the methodology provided a solid basis for understanding artificial intelligence as a contemporary socio-legal phenomenon.

THEORETICAL FRAMEWORK

Artificial intelligence has evolved from a narrow computational discipline into a multidimensional concept studied across sociology, law, economics, and philosophy. Theoretical discussions commonly define AI as the capacity of machines to execute tasks associated with human cognition, such as learning, pattern recognition, language processing, and autonomous decision-making. However, beyond technical definitions, contemporary scholarship interprets AI as a socio-technical system embedded in institutions, norms, and power relations. Floridi and Cowls (2022) argued that AI should be understood through an ethical framework that combines innovation with principles such as autonomy, justice, beneficence, and explicability. This perspective emphasizes that AI systems are never neutral because they reflect human choices, data structures, and governance priorities. Therefore, any theoretical examination of AI must integrate both technological and societal dimensions.

One foundational concept in the literature is machine learning, which refers to computational systems capable of improving performance through exposure to data rather than explicit programming. Alpaydin (2021) explained that machine learning allows systems to adapt based on experience, making it central to modern AI applications. Through supervised, unsupervised, and reinforcement learning methods, machines can identify patterns in vast datasets that exceed human analytical capacity. This capability has transformed fields such as healthcare diagnostics, logistics, finance, and public administration. Yet the dependence on historical data also introduces risks because flawed or biased datasets can generate distorted outcomes. Thus, machine learning is simultaneously a source of innovation and a mechanism through which past inequalities may be reproduced.

Deep learning represents a more advanced branch of machine learning based on artificial neural networks inspired by the human brain. Mehrish et al. (2024) noted that deep learning systems excel in image recognition, speech analysis, translation, and predictive modeling because they can process multiple layers of abstraction. This technological leap has enabled applications ranging from medical imaging to autonomous driving. Nevertheless, many deep learning models function as "black boxes," meaning their internal reasoning is difficult for users or regulators to interpret. Theoretical debates therefore focus not only on performance but also on transparency and trustworthiness. As AI becomes more influential in high-stakes decisions, explainability becomes a necessary condition for legitimate deployment.

Natural language processing (NLP) is another key area of AI theory because language remains one of the most complex forms of human intelligence. Hickok (2022) explained that NLP seeks to enable machines to interpret, generate, and respond to human language in meaningful ways. Applications include chatbots, translation software, automated summarization, sentiment analysis, and legal document review. In sociological terms, NLP changes communication by mediating interactions between individuals, organizations, and digital platforms. It also shapes access to information because algorithms prioritize certain

meanings, topics, or narratives over others. Consequently, language technologies influence public discourse and cultural production in ways that extend beyond efficiency.

From a sociological standpoint, technological determinism has often been used to explain the transformative role of AI. This theory proposes that technological innovation acts as a primary driver of social change. While critics consider this approach overly simplistic, Smith and Marx (1994) observed that technologies frequently reorganize institutions, labor systems, and daily routines. AI illustrates this process through automation, platform economies, and algorithmic governance. However, social constructionist scholars argue that society also shapes technology through regulation, market incentives, and cultural expectations. Therefore, AI should be viewed not as an autonomous force but as a product of reciprocal interaction between technical design and social context.

Another relevant theoretical lens is Pierre Bourdieu's concept of capital and social reproduction. Verwiebe and Hagemann (2025) identified economic, cultural, and social capital as resources that structure inequality across generations. In the AI era, digital skills, data literacy, and technological access function as new forms of capital. Individuals and institutions possessing these resources gain competitive advantages in education, employment, and political influence. Conversely, communities lacking digital infrastructure or advanced training risk exclusion from emerging opportunities. This interpretation helps explain why AI can widen inequality even while promising efficiency and growth. Access to technology alone is insufficient unless accompanied by inclusive educational and social policies.

The digital divide remains central to understanding AI diffusion. Van Deursen and Van Dijk (2019) argued that inequality in the digital age extends beyond internet connectivity to include disparities in skills, motivation, and usage opportunities. AI intensifies this divide because advanced systems often require high-quality infrastructure, expensive computing resources, and specialized expertise. Countries or populations with limited digital capacity may become dependent on technologies developed elsewhere, reinforcing global asymmetries. At the domestic level, rural populations, older adults, and low-income groups may be excluded from AI-enabled services. Hence, the theoretical concept of digital inequality is essential for evaluating the distributive effects of AI adoption.

Labor market transformation is another major area of scholarship. Howcroft and Taylor (2023) explained that automation does not simply eliminate jobs; rather, it restructures tasks within occupations, replacing some functions while complementing others. AI extends automation into professional and analytical domains once considered secure from technological substitution. Routine legal drafting, customer service, accounting, and diagnostics can now be partially automated. At the same time, demand increases for roles involving creativity, emotional intelligence, and system oversight. This task-based framework is important because it moves beyond simplistic narratives of total job loss and instead highlights dynamic reconfiguration of work.

The concept of platform capitalism also contributes to theoretical understanding of AI. Srnicek (2021) described digital platforms as business models that accumulate value through data extraction, network effects, and algorithmic coordination. Major technology firms use AI to optimize advertising, personalize content, manage labor, and predict consumer behavior. In this model, data becomes a strategic asset, while users generate economic value through their interactions. Such concentration of informational power raises concerns about monopolies, democratic accountability, and labor precarity. Therefore, AI must be studied not only as technology but also as an economic infrastructure of contemporary capitalism.

Surveillance theory is equally relevant in the age of AI. Lyon (2019) argued that digital surveillance has moved beyond traditional state observation toward continuous monitoring embedded in everyday systems. AI-powered facial recognition, predictive policing, biometric authentication, and behavioral analytics exemplify this shift. These tools can improve efficiency and security, yet they also normalize constant observation and asymmetrical power

relations. When surveillance disproportionately targets marginalized groups, it can reinforce discrimination under a veneer of technological objectivity. Theoretical engagement with surveillance is thus necessary to assess the balance between security and civil liberty.

From a legal-theoretical perspective, accountability is one of the most debated issues surrounding AI. Jamil et al. (2024) emphasized that legal responsibility traditionally depends on agency, intention, and rule-governed conduct. AI systems complicate these assumptions because they may generate outcomes not directly anticipated by designers or operators. For example, autonomous vehicles or algorithmic decision systems can cause harm without clear human negligence. This creates challenges for tort law, administrative law, and criminal responsibility. Legal theory must therefore reconsider how causation and fault are assigned in contexts involving distributed human-machine interaction.

Data protection theory also provides a crucial framework for AI governance. Solove (2010) argued that privacy should not be reduced to secrecy alone but understood as a set of problems involving information collection, aggregation, exclusion, and misuse. AI systems rely heavily on data collection to train models and personalize outputs. Without safeguards, individuals may lose control over how their personal information is processed, monetized, or shared. This concern has motivated regulatory responses such as consent requirements, access rights, and algorithmic transparency measures. Consequently, privacy theory remains central to balancing innovation with personal autonomy.

Ethics of fairness and non-discrimination have become especially important in AI research. Ferrara et al. (2024) demonstrated that algorithmic systems may produce discriminatory effects even when no explicit bias is intended. Variables correlated with race, gender, or socioeconomic status can reproduce historical disadvantage through automated scoring systems. Fairness theory therefore examines how to detect, measure, and mitigate unequal outcomes. It also highlights that technical solutions alone are insufficient because fairness involves moral and political judgments. Effective governance requires interdisciplinary collaboration among engineers, lawyers, and social scientists.

Finally, human-centered AI has emerged as an integrative theoretical model for the future. Shneiderman (2022) proposed that trustworthy AI should prioritize human values, democratic oversight, safety, and user empowerment rather than unrestricted automation. This approach rejects the idea that efficiency is the sole criterion of success. Instead, it calls for systems that enhance human capabilities while preserving dignity, rights, and social inclusion. Human-centered AI synthesizes sociological concerns about inequality with legal concerns about accountability and rights protection. For that reason, it offers one of the most comprehensive frameworks for guiding responsible technological development.

DISCUSSION

The findings of this study confirm that artificial intelligence has moved beyond its original technical boundaries and now functions as a structural force influencing social organization, economic systems, and legal institutions. As Yadav et al. (2020) explained, AI was initially conceived as a field focused on replicating intelligent behavior through machines; however, its present significance lies in how these systems shape decision-making processes in everyday life. This transformation supports the argument that AI must be analyzed not only through engineering perspectives but also through interdisciplinary frameworks capable of addressing its broader consequences. In this sense, the results align with Floridi and Cowls (2022), who emphasized that AI governance requires ethical principles integrated with innovation.

From a sociological perspective, the discussion reveals that AI simultaneously creates opportunities and reproduces inequalities. On one hand, intelligent systems improve productivity, facilitate access to services, and optimize organizational efficiency. On the other hand, their benefits are distributed unevenly across social groups. Brynjolfsson and McAfee

(2014) argued that automation increases productivity while displacing workers in vulnerable occupations, a conclusion reinforced by Howcroft and Taylor (2023), who observed that technology transforms tasks rather than simply eliminating jobs. The present analysis suggests that AI intensifies this restructuring by extending automation into cognitive and professional sectors. Therefore, labor policies, retraining programs, and inclusive educational strategies are necessary to prevent deeper socioeconomic fragmentation.

The evidence also demonstrates that digital inequality remains one of the central risks associated with AI expansion. Access to infrastructure alone does not guarantee meaningful participation in technologically advanced societies. Van Deursen and Van Dijk (2019) noted that digital divides include disparities in skills, motivation, and effective use of technologies. This perspective is especially relevant in developing contexts where limited connectivity, weak institutional capacity, or insufficient training may hinder AI adoption. Consequently, nations and communities lacking technological capital risk becoming dependent consumers rather than active producers of innovation. This reinforces Verwiebe and Hagemann (2025) theory that new forms of capital can reproduce existing hierarchies unless redistributive mechanisms are implemented.

Another important discussion point concerns surveillance and algorithmic control. AI-powered systems increasingly monitor behavior through facial recognition, predictive analytics, and personalized content curation. Lyon (2019) warned that surveillance has become embedded in routine digital environments rather than confined to exceptional state practices. Likewise, Zuboff (2023) described how data extraction models convert personal behavior into economic value. These concerns are consistent with the findings of this study, which indicate that AI can strengthen asymmetrical power relations between citizens, corporations, and governments. As a result, the challenge is not merely technical efficiency, but the preservation of autonomy, privacy, and democratic accountability.

Algorithmic bias constitutes another major issue requiring critical reflection. Although AI systems are often perceived as objective, they frequently inherit distortions from the data used to train them. O'Neil (2016) demonstrated that automated models can reproduce discrimination in hiring, policing, education, and finance. Similarly, Ferrara et al. (2024) showed that fairness in machine learning cannot be guaranteed solely through technical optimization because normative judgments are always involved. The findings of this research support the view that bias mitigation requires multidisciplinary collaboration, combining statistical methods with legal safeguards and social oversight. Neutrality cannot be assumed simply because a decision is automated.

From a legal standpoint, the results highlight persistent uncertainty regarding liability and regulation. Traditional legal systems were designed around human intention and identifiable responsibility, yet AI systems may generate harmful outcomes through distributed interactions among developers, operators, and autonomous processes. Abbott (2020) emphasized that such scenarios create new dilemmas in tort law and product liability. Jamil et al. (2024) theory of responsibility becomes more difficult to apply when causation is technologically mediated. This study therefore supports the argument that adaptive legal frameworks are necessary to clarify accountability without obstructing beneficial innovation.

Privacy and data governance also remain central to the debate. Solove (2010) argued that privacy problems involve collection, aggregation, exclusion, and misuse of personal information rather than secrecy alone. Since AI depends heavily on data-driven learning, weak safeguards may expose individuals to profiling, manipulation, or unauthorized surveillance. Wachter et al. (2017) further emphasized the importance of explainability and rights related to automated decisions. In practical terms, citizens increasingly demand transparency regarding how algorithms evaluate them and how data is processed. Thus, future regulation must guarantee informed consent, oversight, and meaningful remedies.

Overall, the discussion confirms that AI should not be understood as an autonomous destiny but as a human-designed system shaped by political, economic, and ethical choices. Shneiderman (2022) proposed that human-centered AI offers a more sustainable path by prioritizing safety, dignity, and democratic values. The present study strongly supports this approach. If governed responsibly, AI can expand welfare, efficiency, and knowledge creation. If left unchecked, it may intensify inequality, surveillance, and legal uncertainty. Therefore, the future of artificial intelligence depends less on computational power than on the collective capacity of societies to guide innovation toward justice and inclusion.

CONCLUSIONS

Artificial intelligence has consolidated itself as one of the most influential forces of contemporary society, transforming economic activities, institutional processes, and interpersonal relations. This study demonstrated that AI should not be interpreted solely as a technological advancement, but rather as a socio-legal phenomenon whose effects extend across labor markets, governance systems, and cultural dynamics. Its growing presence in decision-making processes confirms that intelligent systems now play a central role in shaping opportunities, risks, and social behavior. Therefore, understanding AI requires an interdisciplinary perspective capable of integrating technical innovation with human values and legal safeguards.

From a sociological perspective, the research concluded that AI offers significant benefits in productivity, efficiency, and access to services, while simultaneously generating risks of exclusion and inequality. Automation may restructure employment by replacing routine tasks and redefining professional functions, especially in sectors dependent on repetitive cognitive work. Likewise, algorithmic systems may reinforce pre-existing disparities when biased data or unequal access to digital resources determines outcomes. These findings indicate that technological progress alone does not guarantee social progress. Public policies centered on education, digital inclusion, and labor adaptation are essential to ensure that AI contributes to collective welfare rather than deepening structural divisions.

From a legal perspective, the study identified persistent challenges concerning accountability, privacy, transparency, and intellectual property. Existing legal frameworks, originally designed for human actors and conventional products, often struggle to address autonomous or semi-autonomous systems capable of producing unforeseen consequences. Questions regarding liability for harmful decisions, protection of personal data, and ownership of AI-generated content remain unresolved in many jurisdictions. Consequently, modern regulatory approaches must evolve toward flexible and preventive models that protect rights without unnecessarily limiting innovation.

In general terms, the future impact of artificial intelligence will depend less on the sophistication of algorithms than on the ethical and institutional choices societies make today. AI can become a powerful instrument for sustainable development, social inclusion, and administrative efficiency if guided responsibly. However, without adequate oversight, it may intensify surveillance, inequality, and legal uncertainty. For this reason, a human-centered model of AI governance is imperative. Balancing innovation with dignity, fairness, and democratic accountability will determine whether artificial intelligence becomes a tool of emancipation or a source of new forms of domination.

LIMITATIONS OF THE STUDY

This study presents certain limitations that should be acknowledged. First, it was developed through a theoretical and documentary approach, which means that the analysis relied primarily on existing academic literature rather than empirical field data. Second, because artificial intelligence evolves rapidly, some technological or regulatory developments may change after the completion of this research. Finally, the socio-legal scope of the study

required the integration of multiple disciplines, which may limit the depth of analysis in some specialized technical areas.

FUTURE STUDIES

Future research should incorporate empirical methodologies, including surveys, interviews, and case studies, to evaluate the real impact of artificial intelligence on institutions, labor markets, and social behavior. Comparative studies between countries would also be valuable to examine differences in regulatory frameworks and public policies. In addition, further research is recommended on emerging issues such as generative AI, algorithmic accountability, digital rights, and the ethical governance of autonomous systems.

ACKNOWLEDGMENTS

The authors express their sincere gratitude to the academic specialists and researchers whose valuable knowledge and scientific contributions supported the development of this study. Special appreciation is extended to professional colleagues and university peers for their constructive observations, intellectual encouragement, and continuous collaboration throughout the research process. The authors also acknowledge the broader academic community whose published works provided an essential foundation for this investigation.

CONFLICT OF INTEREST STATEMENT

The authors declare that there are no conflicts of interest regarding the publication of this study. No financial, institutional, or personal relationships influenced the design, analysis, interpretation, or presentation of the research findings.

REFERENCES

- Abbott, R. (2020). *The Reasonable Robot: Artificial Intelligence and the Law*. Cambridge University Press.
- Alpaydin, E. (2021). *Machine learning*. MIT press.
- Brynjolfsson, E., & McAfee, A. (2014). *The Second Machine Age*. W. W. Norton & Company.
- Castells, M. (2010). *The Rise of the Network Society* (2nd ed.). Wiley-Blackwell.
- Espinoza Freire, E. E. (2020). La búsqueda de información científica en las bases de datos académicas. *Revista Metropolitana de Ciencias Aplicadas*, 3(1), 31-35.
- Espinoza Freire, E. E. (2020). La investigación cualitativa, una herramienta ética en el ámbito pedagógico. *Conrado*, 16(75), 103-110.
- Espinoza Freire, E. E., & Calva Nagua, D. X. (2020). La ética en las investigaciones educativas. *Revista Universidad y sociedad*, 12(4), 333-340.
- Espinoza Freire, E. E., & Rad Camayd, Y. (2020). A ética na pesquisa inclusiva, uma ferramenta didáctica. *Revista Universidad y Sociedad*, 12(6), 139-146.
- Espinoza-Freire, E. E. (2021). La argumentación científica una herramienta didáctica. *Revista Uniandes Episteme*, 8(1), 106-121.
- Espinoza-Freire, E. E. (2022). Ética en la investigación científica. *Revista Mexicana de Investigación e Intervención Educativa*, 1(2), 35-43.
- Espinoza-Freire, E. E. (2024). Teoría Fundamentada: evolución, principios y aplicaciones en la investigación cualitativa. *Sophia Research Review*, 1(1), 26-33
- Ferrara, C., Sellitto, G., Ferrucci, F., Palomba, F., & De Lucia, A. (2024). Fairness-aware machine learning engineering: how far are we?. *Empirical software engineering*, 29(1), 9.

- Floridi, L., & Cows, J. (2022). A unified framework of five principles for AI in society. *Machine learning and the city: Applications in architecture and urban design*, 535-545.
- Hickok, G. (2022). The dual stream model of speech and language processing. *Handbook of clinical neurology*, 185, 57-69.
- Howcroft, D., & Taylor, P. (2023). Automation and the future of work: A social shaping of technology approach. *New technology, work and employment*, 38(2), 351-370.
- Jamil, J., Fadli, M., Hadiyantina, S., & Prasetyo, N. D. (2024). Redesigning the Concept of Law Enforcement in Administrative Violations of General Elections in Indonesia. *YURIDIKA*, 39(3), 279-302.
- Lyon, D. (2019). Surveillance capitalism, surveillance culture and data politics 1. In *Data politics* (pp. 64-77). Routledge.
- McCarthy, J., Minsky, M., Rochester, N., & Shannon, C. (2006). A proposal for the Dartmouth summer research project on artificial intelligence, August 31, 1955. *AI Magazine*, 27(4), 12-14.
- Mehrish, A., Majumder, N., Bharadwaj, R., Mihalcea, R., & Poria, S. (2023). A review of deep learning techniques for speech processing. *Information Fusion*, 99, 101869.
- O'Neil, C. (2016). *Weapons of Math Destruction*. Crown Publishing.
- Shneiderman, B. (2022). *Human-Centered AI*. Oxford University Press.
- Smith, M. R., & Marx, L. (Eds.). (1994). *Does technology drive history?: The dilemma of technological determinism*. Mit Press.
- Solove, D. J. (2010). *Understanding Privacy*. Harvard University Press.
- Srnicek, N. (2021). Value, rent and platform capitalism. In *Work and labour relations in global platform capitalism* (pp. 29-45). Edward Elgar Publishing.
- Van Deursen, A. J., & Van Dijk, J. A. (2019). The first-level digital divide shifts from inequalities in physical access to inequalities in material access. *New media & society*, 21(2), 354-375.
- Verwiebe, R., & Hagemann, S. (2025). Bourdieu revisited: new forms of digital capital-emergence, reproduction, inequality of distribution. *Information, Communication & Society*, 28(11), 1861-1883.
- Wachter, S., Mittelstadt, B., & Floridi, L. (2017). Why a right to explanation of automated decision-making does not exist in the General Data Protection Regulation. *International Data Privacy Law*, 7(2), 76-99.
- Yadav, S. P., Mahato, D. P., & Linh, N. T. D. (Eds.). (2020). *Distributed artificial intelligence: A modern approach*. CRC Press.
- Zuboff, S. (2023). The age of surveillance capitalism. In *Social theory re-wired* (pp. 203-213). Routledge.