

RESEARCH ESSAY

Human dignity in the age of artificial intelligence, legal and technological challenges of the twenty-first century

Dignidad humana en la era de la inteligencia artificial desafíos jurídicos y tecnológicos del siglo XXI

Javier A. Artiles-Santana 

Received: 19 June 2025 / Accepted: 11 July 2025 / Published online: 5 August 2026
© The Author(s) 2026

Abstract This study critically examines the contemporary reconfiguration of the legal understanding of human dignity in the digital era and the legal and technological challenges arising from the increasing incorporation of artificial intelligence systems into decision-making processes that may affect fundamental rights. The analysis draws primarily on scientific literature retrieved from Scopus and Web of Science, complemented by doctrinal, normative, regulatory, and institutional sources issued by international organizations. A qualitative documentary design was adopted, combining legal and philosophical analysis within the frameworks of legal constructivism and the paradigm of complexity. The results reveal tensions involving human autonomy, privacy, equality, non-discrimination, transparency, accountability, and effective human oversight. From the documentary analysis, the study develops the construct of systemic dignity, understood as the dynamic projection of human dignity within complex sociotechnical environments in which individuals, legal systems, and artificial intelligence technologies interact. This construct extends the explanatory and regulatory scope of human dignity and supports rights-based, adaptive, multilevel regulatory frameworks that ensure effective human control and concrete legal safeguards throughout the artificial intelligence life cycle. It thereby reinforces the primacy of the person and protection of fundamental rights in automated decision-making environments.

Keywords: artificial intelligence; human dignity; fundamental rights; algorithmic regulation; systemic dignity.

Resumen Este estudio examina críticamente la comprensión jurídica de la dignidad humana en la era digital y los desafíos jurídicos y tecnológicos derivados de la incorporación de sistemas de inteligencia artificial en procesos de decisión que pueden afectar derechos fundamentales. El análisis se basa en literatura científica de Scopus y Web of Science, complementada con fuentes doctrinales, normativas, regulatorias e institucionales de organismos internacionales. Se adoptó un diseño documental cualitativo que combina el análisis jurídico y filosófico desde el constructivismo jurídico y el paradigma de la complejidad. Los resultados revelan tensiones en torno a la autonomía humana, la privacidad, la igualdad, la no discriminación, la transparencia, la rendición de cuentas y la supervisión humana efectiva. A partir del análisis documental, el estudio desarrolla el constructo de dignidad sistémica, entendido como la proyección dinámica de la dignidad humana en entornos sociotécnicos complejos donde interactúan personas, sistemas jurídicos y tecnologías de inteligencia artificial. Este constructo amplía el alcance explicativo y regulatorio de la dignidad humana y respalda marcos regulatorios basados en derechos, adaptativos y multinivel, capaces de garantizar control humano efectivo y salvaguardias jurídicas durante el ciclo de vida de la inteligencia artificial.

Palabras clave : inteligencia artificial; dignidad humana; derechos fundamentales; regulación algorítmica; dignidad sistémica.

How to cite

Artiles-Santana, J. A. (2026). Human dignity in the age of artificial intelligence, legal and technological challenges of the twenty-first century. Journal of Law and Epistemic Studies, 4, e195. <https://doi.org/10.5281/zenodo.22331033>

✉ Javier A. Artiles-Santana
jartilessantana1@gmail.com
Universidad San Gregorio de Portoviejo,
Manabí, Ecuador.

PUBLIS
EDITORIAL



Introduction

Artificial intelligence (AI) constitutes one of the most significant technological transformations of the contemporary era. Its modern development is commonly associated with Alan Turing's early reflections on machine intelligence and with the Dartmouth Conference of 1956, widely regarded as a foundational moment in the consolidation of artificial intelligence as a distinct field of research (Turing, 2009; Arrestegui, 2012).

Since then, AI has progressively transformed numerous areas of social, economic, and institutional life, including automation, healthcare, transportation, communications, education, scientific research, and large-scale data processing. These developments have contributed to greater efficiency, expanded analytical capacity, and the optimization of multiple human activities. At the same time, they have generated increasingly complex legal and ethical concerns regarding the possible impact of automated systems on fundamental values and legally protected interests, particularly human dignity (Russell & Norvig, 2016; Brundage, 2015).

The rapid expansion of generative artificial intelligence has intensified this debate. AI systems are now being incorporated into legally and socially sensitive areas such as the administration of justice, healthcare, public security, employment, education, and public services, where automated recommendations, classifications, or decisions may have direct consequences for individuals. This development has strengthened the need to identify legal and ethical principles capable of reconciling technological innovation with the effective protection of human rights and fundamental freedoms (Floridi et al., 2018; Sun, 2018; Belouali et al., 2020; Ruttkamp-Bloem, 2020; Floridi & Cowls, 2021; Akula & Garibay, 2021; Wong, 2021; Heiling, 2022; Hagendorff, 2022).

Recent scholarship confirms that the legal implications of artificial intelligence are not confined to isolated infringements of specific rights. Teo (2025) warns that the gradual incorporation of AI into social and institutional decision-making may affect the normative foundations of human-rights protection itself. This perspective highlights the need to analyze the relationship between artificial intelligence and human dignity not only in terms of specific technological harms, but also in relation to the broader legal architecture through which human rights are interpreted, guaranteed, and enforced.

International organizations have progressively incorporated these concerns into ethical and regulatory instruments. UNESCO's Recommendation on the Ethics of Artificial Intelligence places human dignity, human rights, fairness, transparency, accountability, and human oversight among the central principles that should govern the development and use of AI systems (UNESCO, 2021).

Similarly, the European Union has developed a progressively binding regulatory framework designed to ensure that artificial intelligence operates within conditions compatible with safety, transparency, non-discrimination, accountability, and the protection of fundamental rights.

Nevertheless, the regulation of artificial intelligence continues to face substantial challenges arising from the complexity of the interaction between technological systems and legally protected rights (Dignum, 2019). Legal and philosophical scholarship has consequently intensified its examination of the relationship among artificial intelligence, justice, human dignity, fundamental rights, and the rule of law (Llano et al., 2022; Consejo General del Poder Judicial, 2024). In this regard, Shaelou and Razmetaeva (2023) emphasize that the emerging digital legal order must remain subject to rule-of-law guarantees and to the effective protection of fundamental rights, particularly where technological systems acquire a significant role in processes of decision-making and governance.

Human dignity has historically occupied a foundational position within contemporary constitutional and international human-rights law. The Universal Declaration of Human Rights establishes that all human beings are born free and equal in dignity and rights, while the International Covenant on Civil and Political Rights and the International Covenant on Economic, Social and Cultural Rights reaffirm the inherent dignity of the human person as a foundational basis of the international human-rights system (United Nations, 1948, 1966a, 1966b).

From a philosophical perspective, Kant's conception of dignity remains particularly influential. Kant (2012) links dignity to the rational and autonomous character of the human person and to the imperative that individuals must always be treated as ends in themselves and never merely as means. Under this approach, dignity constitutes an intrinsic value without equivalent and serves as a normative barrier against the instrumentalization of the person. This understanding has profoundly influenced contemporary constitutional theory and the interpretation of dignity as a moral and legal foundation of fundamental rights (Habermas, 2010).

The emergence of increasingly sophisticated artificial systems has renewed debate regarding the meaning and scope of dignity in technological environments. Thielscher (2025) argues that the development of artificial intelligence requires greater conceptual precision concerning dignity within computer ethics, particularly in relation to freedom, autonomy, moral responsibility, and individuality. Although this debate also extends to questions concerning the possible moral status of machines, the present study remains grounded in a human-centered conception of dignity and focuses on the legal implications generated when automated systems intervene in decisions affecting individuals.

This intervention creates important legal tensions. Decisions traditionally dependent on human judgment may now be adopted, recommended, prioritized, or conditioned by algorithmic systems. Such processes may generate opacity, discriminatory outcomes, profiling, surveillance, and difficulties in attributing legal responsibility, particularly where automated decisions affect access to employment, social benefits, credit, public services, justice, education, or other legally protected opportunities and rights (Müller, 2023).

Bryson (2016) acknowledges that artificial intelligence may contribute significantly to human welfare by improving services and strengthening decision-making through the analysis of large volumes of information. However, its use also raises fundamental questions regarding autonomy, responsibility, and the conditions under which technological systems may influence human conduct. Bostrom and Yudkowsky (2018) similarly examine the ethical consequences associated with highly advanced artificial intelligence systems and the potential risks arising from insufficient human control.

Coeckelbergh (2020) further warns that technological development without adequate legal and ethical safeguards may contribute to processes of dehumanization and weaken human autonomy. These positions show that the central legal issue is not whether artificial intelligence is beneficial or harmful in abstract terms, but under what normative conditions its design and use remain compatible with the protection of the person as a bearer of dignity and fundamental rights.

The relationship between artificial intelligence and law is also characterized by a significant asymmetry in their respective rates of development. Technological innovation evolves rapidly and continuously, whereas legal systems require comparatively slower processes of legislative adoption, judicial interpretation, administrative implementation, and institutional adaptation. This temporal asymmetry creates difficulties for the timely construction of regulatory safeguards capable of responding to the risks generated by automated decision-making.

Yeung (2018) uses the concept of algorithmic regulation to describe forms of governance increasingly mediated by automated systems and warns that such arrangements may affect fairness, accountability, procedural guarantees, and individual rights. The expansion of algorithmic governance therefore requires a deeper legal analysis of the relationship among technology, power, regulation, and the effective protection of fundamental rights.

This concern is also reflected in recent approaches to algorithmic governance and People Analytics, which emphasize the need to embed transparency, accountability, non-discrimination, and labour-rights safeguards into digitally mediated organizational decision-making (Esquivel García & Vargas Mursulí, 2026).

The recent literature reinforces this concern. Orwat et al. (2024) argue that risk-based regulation of artificial intelligence faces particular difficulties when legal systems must operationalize normative principles such as dignity, informational self-determination, privacy, non-discrimination, fairness, and justice. This difficulty demonstrates that the challenge is not limited to identifying abstract values, but extends to translating those values into concrete legal obligations, compliance mechanisms, procedural safeguards, and review structures.

The European Union's regulatory response illustrates this shift from general principles toward enforceable legal obligations. Regulation (EU) 2024/1689 establishes a risk-based regulatory architecture for artificial intelligence and assigns particular relevance to the protection of fundamental rights. In the case of high-risk AI systems, the Regulation requires measures of human oversight designed to prevent or minimize risks to health, safety, and fundamental rights, thereby confirming that effective human control constitutes a core legal safeguard in automated environments (European Parliament & Council of the European Union, 2024).

Despite these advances, the global regulatory landscape remains fragmented. Significant differences persist among jurisdictions regarding the scope, intensity, and enforceability of AI regulation. In Latin America, for example, regulatory initiatives continue to display heterogeneous levels of development and institutional consolidation (Hernández et al., 2024). This fragmentation reinforces the existence of a normative gap between rapid technological innovation and the capacity of legal systems to provide coherent, effective, and internationally compatible safeguards.

The problem therefore requires an analytical framework capable of integrating legal, philosophical, technological, and ethical dimensions. Legal constructivism provides a useful theoretical basis because it recognizes that legal concepts are not immutable, but evolve as social structures, institutional practices, and normative expectations change (Dworkin, 1986; Cáceres, 2007).

From this perspective, human dignity may be understood as a legal construct whose protective function must be capable of responding to new forms of social interaction mediated by artificial intelligence. This does not imply abandoning its classical foundations. Rather, it requires examining how dignity can continue to operate as a normative limit and interpretative criterion in contexts where algorithmic systems increasingly participate in decisions with legal and social consequences.

The paradigm of complexity complements this approach. Morin (2002) proposes a dialogical model that makes it possible to analyze apparently opposing dimensions as interconnected components of the same complex reality. Applied to the present study, this requires moving beyond a fragmented analysis of artificial intelligence and human

rights and examining instead the legal, ethical, technological, and institutional relations that connect them.

Within this system of relations, human dignity functions as the principal normative point of articulation between technological innovation and the protection of fundamental rights. Artificial intelligence systems operate within complex sociotechnical structures in which legal norms, ethical principles, institutional practices, technological architecture, and social consequences interact continuously (Floridi, 2013). Consequently, AI governance cannot be addressed exclusively through isolated technological or legal measures.

The analysis of these relationships makes it possible to identify four critical areas in which artificial intelligence may generate particularly significant tensions with human dignity: automated decision-making and algorithmic autonomy; privacy and surveillance; equality and non-discrimination; and transparency and accountability in governance. In these areas, automated systems may reinforce structural inequalities, weaken individual autonomy, obstruct effective review, or affect access to rights and opportunities when adequate legal safeguards are absent (Eubanks, 2018).

Recent research on ethics-by-design further strengthens the need to incorporate rights-based safeguards at the earliest stages of technological development. Brey and Dainow (2024) propose integrating human agency, fairness, transparency, privacy, and responsibility into the design and development of artificial intelligence systems. This approach is consistent with the view that human dignity should not function merely as a remedial principle invoked after harm has occurred, but also as a preventive legal criterion capable of guiding technological design before risks materialize.

On this basis, the study develops the construct of systemic dignity. Systemic dignity is defined as the dynamic projection of human dignity within complex sociotechnical environments, intended to ensure that the design, deployment, use, and supervision of artificial intelligence systems remain subordinated to human autonomy, fundamental rights, legal accountability, transparency, and effective human control.

This construct does not replace the classical conception of dignity nor does it modify its fundamentally human character. Rather, it extends its explanatory and regulatory capacity to technological environments in which the exercise of rights, access to opportunities, and institutional decision-making are increasingly mediated by automated systems.

Under this approach, dignity acquires not only a foundational but also an operative function. It may serve as a legal criterion for the design, implementation, assessment, supervision, and review of artificial intelligence systems, particularly where such systems affect fundamental rights or legally protected interests.

The research is therefore guided by the following question:

How does artificial intelligence affect human dignity within contemporary legal systems, and how can systemic dignity contribute to strengthening legal protection against the risks associated with automated decision-making?

Accordingly, the general objective of the study is to critically analyze the relationship between human dignity and artificial intelligence from legal and philosophical perspectives and to develop systemic dignity as a conceptual and normative category capable of strengthening the protection of fundamental rights in automated environments.

The relevance of the study lies in the need to develop legal responses capable of reconciling technological innovation with the protection of the person. In this context, this framework is proposed as a theoretical and normative framework that may contribute to the interpretation, regulation, and supervision of artificial intelligence systems where automated processes affect autonomy, freedom, equality, security, privacy, and other fundamental rights.

Methodology

The research was developed through a qualitative documentary design, appropriate for examining the legal and philosophical implications arising from the interaction between human dignity and artificial intelligence systems. This methodological approach made it possible to analyze the problem from a normative, doctrinal, conceptual, and interdisciplinary perspective, consistent with the objective of identifying legal tensions and formulating a theoretical construct capable of responding to emerging technological realities.

The principal research technique was documentary review and legal-documentary analysis. The documentary corpus was composed primarily of scientific publications retrieved from the Scopus and Web of Science databases, together with doctrinal, normative, regulatory, and institutional sources directly related to artificial intelligence, human dignity, human rights, algorithmic regulation, and automated decision-making systems. Institutional documents issued by international organizations, particularly UNESCO and the European Union, were also examined because of their relevance to the ethical governance and legal regulation of artificial intelligence.

The selection of documents was guided by criteria of thematic relevance, legal significance, philosophical pertinence, and direct connection with the object of study. Priority was given to sources addressing the protection of fundamental rights in technological environments, the ethical and legal implications of artificial intelligence, algorithmic governance, automated decision-making, human oversight, transparency, non-discrimination, privacy,

accountability, and regulatory responses to technological innovation.

The documentary review was oriented toward identifying conceptual convergences, doctrinal divergences, normative tensions, regulatory gaps, and emerging legal criteria capable of explaining the relationship between artificial intelligence and human dignity. The analysis therefore went beyond a merely descriptive review of the literature and sought to determine how the progressive incorporation of automated systems into legally sensitive contexts may affect the exercise and protection of fundamental rights.

The study combined a legal and philosophical analytical perspective. From the legal standpoint, the research examined rules, principles, regulatory instruments, institutional standards, and governance mechanisms concerning the protection of human dignity and fundamental rights in the context of artificial intelligence. Particular attention was given to legal safeguards related to transparency, accountability, non-discrimination, privacy, human oversight, and the possibility of reviewing or contesting decisions produced or influenced by automated systems.

From the philosophical perspective, the research examined the classical foundations of human dignity and their capacity to respond to contemporary technological environments. This dimension made it possible to assess whether the traditional understanding of dignity remains sufficient in situations where algorithmic systems participate in processes capable of affecting autonomy, freedom, equality, privacy, and other legally protected interests.

The theoretical and methodological framework integrates legal constructivism with the paradigm of complexity. Legal constructivism provided the basis for understanding legal concepts as interpretative categories capable of evolving in response to transformations in social, institutional, and technological realities. Under this approach, human dignity was not treated as an immutable or closed category, but as a foundational legal principle whose protective function may require renewed forms of interpretation when new relationships emerge between individuals and automated systems.

The paradigm of complexity, in turn, made it possible to analyze the relationship among artificial intelligence, law, ethics, philosophy, technology, and society as components of a single interconnected system. This perspective avoided treating technological development and fundamental-rights protection as isolated dimensions and instead examined the reciprocal influences that arise between them.

Within this analytical framework, human dignity was considered the principal normative point of articulation between artificial intelligence systems and fundamental rights. This approach enabled the study to examine both directions of the relationship: how technological systems may affect legally protected rights and interests, and how

legal principles and human-rights standards may condition the design, deployment, use, and supervision of artificial intelligence.

Given the novel and evolving character of the object of study, the documentary analysis followed an interpretative and inductive procedure. Through this process, recurring legal categories, normative tensions, regulatory deficiencies, and patterns of convergence were identified across the scientific, doctrinal, and institutional sources reviewed.

The analysis led to the organization of the findings into four principal dimensions:

Conceptual dimension, focused on the contemporary meaning and scope of human dignity in automated environments. Normative dimension, centered on the regulatory gap between technological development and the legal mechanisms intended to protect fundamental rights. Systemic dimension, concerned with the interaction among artificial intelligence, law, ethics, human rights, and technological infrastructures. Prospective dimension, oriented toward governance mechanisms, prevention, human oversight, legal accountability, and adaptive regulation. These analytical dimensions structure the Results and Discussion section and provide the conceptual basis for the formulation of the construct of systemic dignity.

The proposed construct was developed through an interpretative synthesis of the documentary findings. It is understood as a legal and philosophical construct, rather than as an empirically measured variable. Its purpose is to articulate the protection of human autonomy, fundamental rights, transparency, accountability, legal control, and human oversight within technological environments increasingly characterized by automated decision-making.

Accordingly, the methodological approach made it possible to analyze the relationship between artificial intelligence and human dignity not merely as a technological issue, but as a complex legal phenomenon requiring interdisciplinary and transdisciplinary interpretation. The combined use of legal constructivism and the complexity paradigm provided the analytical basis for understanding dignity as both a foundational principle and a potentially operative criterion for the regulation and supervision of artificial intelligence systems.

Results and Discussion

The documentary analysis reveals significant tensions arising from the interaction between artificial intelligence systems and the protection of human rights, with human dignity occupying a central position within that relationship. The findings are organized into four analytical dimensions: conceptual, normative, systemic, and prospective. This structure makes it possible to distinguish the theoretical implications of AI for the notion of dignity, the regulatory limitations of existing legal frameworks, the systemic

interdependence between technology and fundamental rights, and the prospective requirements for future governance and legal safeguards.

The documentary evidence indicates that the classical Kantian conception of human dignity, traditionally understood as an intrinsic and inalienable value inherent in every person, is confronted with new challenges in the context of artificial intelligence. The traditional notion of dignity, grounded in human autonomy and rationality and in the requirement that every person be treated as an end in themselves rather than merely as a means, acquires renewed legal relevance when decisions capable of significantly affecting individuals are adopted, recommended, prioritized, or conditioned by automated systems (Kant, 2012).

This problem is particularly relevant in legally sensitive domains such as the administration of justice, public security, employment, access to social benefits, healthcare, education, financial services, and social inclusion. In these contexts, the partial displacement of human judgment by automated or algorithmically assisted decision-making may alter the traditional structure through which legal responsibility, individual autonomy, procedural guarantees, and the protection of fundamental rights are exercised.

The legal concern therefore does not arise solely from the existence of automated decision-making as a technological phenomenon. It arises from the possibility that algorithmic systems may influence or determine outcomes affecting rights and legally protected interests without sufficient transparency, accountability, contestability, or meaningful human intervention.

Müller (2023) identifies responsibility, autonomy, and control as central ethical and legal issues associated with algorithmic decision-making. In a similar direction, Coeckelbergh (2020) warns that certain applications of artificial intelligence may contribute to processes of dehumanization and may undermine individual autonomy when technological mediation displaces meaningful human judgment. These concerns become particularly relevant where affected persons are unable to understand, challenge, or obtain effective review of decisions produced or substantially influenced by automated systems.

Recent scholarship further demonstrates that the legal implications of artificial intelligence may extend beyond isolated infringements of particular rights. Teo (2025) argues that the effects of AI may gradually reach the normative foundations of the human-rights system itself. This perspective is especially important because it shifts the analysis from individual technological harms toward the broader question of whether automated systems may alter the institutional and normative conditions under which rights are interpreted, exercised, and protected.

The contemporary debate also confirms that increasingly sophisticated artificial systems have renewed discussion concerning the content and scope of human dignity. Thielscher (2025) identifies freedom, autonomy, moral

responsibility, and individuality among the elements that must be taken into account when examining dignity in the field of computer ethics. Although his analysis also addresses the broader philosophical question of whether artificial systems could eventually acquire some form of moral consideration, the present research remains anchored in the legal protection of human dignity and in the consequences that technological systems may generate for persons as holders of fundamental rights.

The findings therefore support the need to reconsider the contemporary legal functionality of human dignity. This does not imply abandoning its classical philosophical foundation, nor does it entail attributing dignity to artificial systems. Rather, it requires determining how human dignity can continue to operate as an effective legal principle in contexts where increasingly complex automated systems mediate decisions affecting individuals.

In this sense, human dignity should not be understood merely as an abstract axiological reference. Within technologically mediated environments, it must also function as a normative limit to the instrumentalization of the person, particularly where algorithmic systems classify, rank, predict, profile, or make recommendations concerning individuals in ways that may affect their rights or life opportunities.

This requires examining dignity not only as a foundational principle of the legal order, but also as a criterion capable of informing the legality, proportionality, legitimacy, and reviewability of automated decision-making processes.

From the interpretative analysis emerges the construct of systemic dignity, understood as a conception of human dignity adapted to the increasingly complex relations among human beings, legal systems, and artificial intelligence technologies.

this framework is conceived as a dynamic and adaptive legal-philosophical category. Its purpose is to ensure that human dignity retains its protective and normative function in environments characterized by continuous technological transformation. The construct responds particularly to the difficulty faced by legal systems in evolving at the same speed as artificial intelligence technologies.

The rapid development of AI may exceed the response capacity of conventional legal frameworks. In this context, the proposed construct seeks to articulate, within the same normative structure, human autonomy, fundamental rights, legal accountability, technological development, and effective human oversight.

Accordingly, systemic dignity does not substitute the traditional concept of human dignity. Rather, it extends its practical and interpretative operation into sociotechnical environments in which technological systems increasingly participate in decisions with legal and social consequences.

From a dialogical perspective, dignity therefore acts as a point of connection between artificial intelligence and the human person. Its function is to permit an integrated interpretation of technological and human dimensions so that the protection of fundamental rights does not remain external to the processes through which AI systems are designed, implemented, used, and supervised.

The complexity paradigm provides an appropriate framework for understanding the interaction between artificial intelligence and human rights as interdependent dimensions rather than as separate systems. Morin (2002) proposes a dialogical approach that allows apparently opposing phenomena to be examined as components of the same complex reality. Applied to the present research, this perspective makes it possible to analyze technological development and the protection of fundamental rights as processes that continuously influence one another.

On one side of this relationship lies the technological system, characterized by automation, large-scale data processing, predictive capacity, and increasingly sophisticated forms of decision support. On the other lies the human and legal system, structured around rights, principles, values, responsibilities, and guarantees. Human dignity operates as the principal normative point of articulation between both dimensions.

The analysis also shows that the relationship between technology and law is characterized by continuous feedback. Technological advances generate new regulatory demands, while legal norms, institutional practices, and social values simultaneously influence the design, deployment, and acceptable use of technological systems. This reciprocal interaction supports the understanding of dignity as a concept that must retain its protective function within continuously changing sociotechnical environments.

Within this framework, the documentary analysis identifies four critical areas in which artificial intelligence may generate particularly significant tensions with human dignity.

The first concerns decision-making and algorithmic autonomy, especially when automated systems intervene in access to resources, opportunities, services, or rights. In these circumstances, the legal concern is whether human agency, meaningful review, and responsibility remain effectively guaranteed.

The second concerns privacy and surveillance, particularly through technological systems capable of collecting, processing, inferring, and using large volumes of personal data. These practices may generate risks to informational self-determination, privacy, and individual autonomy, particularly where surveillance technologies operate without sufficient legal safeguards.

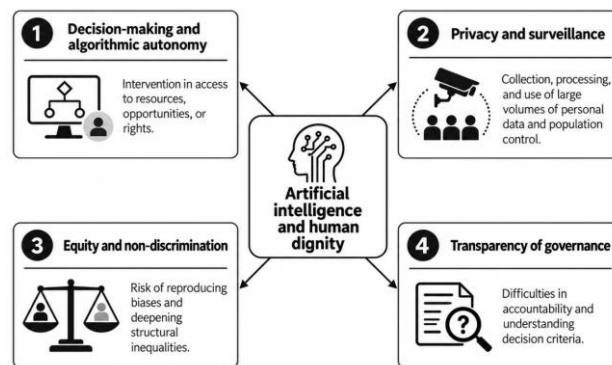
The third concerns equity and non-discrimination, due to the possibility that algorithmic systems may reproduce historical biases, rely on unrepresentative data, or deepen

structural inequalities through classifications or predictions that appear technically neutral.

The fourth concerns transparency of governance, particularly when individuals are unable to understand the criteria underlying automated decisions, identify the responsible actors, or effectively challenge outcomes capable of affecting their rights or legally protected interests.

Figure 1

Critical areas of tension between artificial intelligence and human dignity



Source: Author's own elaboration based on the documentary analysis.

These four areas are consistent with the concerns identified by Eubanks (2018) regarding the automation of inequality and with the normative challenges examined by Orwat et al. (2024), who identify human dignity, privacy, informational self-determination, non-discrimination, equity, and justice among the principal values implicated in AI risk regulation.

This convergence reinforces the systemic approach adopted in the study. The potential effects generated by artificial intelligence do not operate in isolation. A single automated decision may simultaneously involve questions of privacy, discrimination, autonomy, transparency, accountability, and access to opportunities.

Consequently, the legal analysis of artificial intelligence requires an integrated perspective capable of recognizing the interaction among these different dimensions. Fragmented regulatory responses may be insufficient where the effects of technological systems extend across several fundamental rights at the same time.

The findings therefore support the need to approach AI regulation from a complex, interdisciplinary, and multilevel perspective, capable of integrating the legal, technological, ethical, institutional, and social dimensions involved in the relationship between artificial intelligence and human beings.

The findings further indicate that the protection of human dignity in AI-mediated environments requires a regulatory approach capable of moving beyond fragmented and merely reactive responses. Given the cross-border nature of many artificial intelligence systems, exclusively national regulatory solutions may prove insufficient where

technological design, data flows, deployment practices, and legal effects transcend territorial boundaries.

A first requirement is the adoption of a preventive approach. The potential impact of artificial intelligence systems on human dignity and fundamental rights should be assessed before deployment, particularly in contexts involving justice, public security, employment, healthcare, education, public services, social protection, and other areas in which automated decisions may produce significant legal consequences.

This preventive orientation is consistent with the ethics-by-design approach proposed by Brey and Dainow (2024), which incorporates human agency, fairness, transparency, privacy, and responsibility into the design and development of artificial intelligence systems. From a legal perspective, this approach supports the idea that the protection of dignity should not operate only as a remedial principle after harm has occurred, but also as an *ex ante* criterion capable of influencing the technological architecture itself.

A second requirement concerns the development of multilevel governance mechanisms. Cooperation among states, international organizations, regulatory authorities, developers, deployers, and other relevant actors may contribute to establishing common standards for the protection of fundamental rights in relation to technologies whose operation and effects transcend national jurisdictions.

This multilevel approach is particularly relevant in light of the regulatory fragmentation identified in the study. The existence of substantially different legal standards across jurisdictions may create uneven levels of protection and may allow technologically similar systems to operate under divergent safeguards. Greater regulatory coordination could therefore contribute to strengthening legal certainty and ensuring more consistent protection of human dignity.

A third requirement is the development of adaptive regulatory frameworks. The rapid evolution of artificial intelligence makes it necessary to adopt legal mechanisms capable of responding to technological change without weakening the protection of fundamental rights. Regulatory adaptability should not be understood as normative flexibility without limits, but as the capacity of legal systems to update supervisory, interpretative, and enforcement mechanisms while preserving non-negotiable standards of human-rights protection.

A fourth requirement is the guarantee of meaningful human oversight and control. Russell (2019) emphasizes the importance of ensuring that artificial intelligence systems remain compatible with human values and under human control. UNESCO (2021) similarly places human agency and oversight among the central elements of responsible AI governance.

The European regulatory framework reinforces this approach. Regulation (EU) 2024/1689 expressly requires human oversight in relation to high-risk AI systems and establishes that such supervision must be appropriate to the

risks, level of autonomy, and context of use of the system (European Parliament & Council of the European Union, 2024). This confirms the increasing juridification of human oversight as a safeguard against excessive or uncontrolled automation.

The need for effective human control becomes particularly important in areas in which automated systems may have substantial consequences for individuals, including the administration of justice, public security, surveillance, employment, access to essential services, healthcare, and other fields involving fundamental rights.

The findings also show that legal responsibility cannot become diluted merely because a decision is technologically mediated. The involvement of artificial intelligence systems does not eliminate the need to identify the persons or institutions responsible for the design, deployment, supervision, and consequences of automated decision-making. Accountability must therefore remain legally attributable and institutionally enforceable.

Likewise, transparency should not be reduced to the mere disclosure that an artificial intelligence system has been used. Effective transparency requires that individuals affected by automated decisions have access, to the extent legally and technically possible, to meaningful information concerning the factors, criteria, or processes that influenced the outcome, particularly where such decisions affect rights or legally protected interests.

The possibility of challenging or reviewing automated decisions is equally important. Human dignity and due-process guarantees require that individuals should not remain subject to consequential automated outcomes without access to effective mechanisms of contestation, rectification, or human reconsideration.

Floridi (2013) provides a broader ethical framework for understanding these issues through the ethics of information, emphasizing the importance of structuring technological environments in ways compatible with the protection of persons. From this perspective, the central problem is not whether technology is inherently beneficial or harmful, but under what legal and ethical conditions it may be legitimately developed and used.

Within this framework, this framework acquires particular relevance as the principal theoretical contribution of the study. Its function is to articulate the technological dimension with the legal protection of the person, allowing human dignity to operate not only as a foundational value but also as a concrete normative criterion throughout the life cycle of artificial intelligence systems.

Accordingly, the proposed construct may guide the assessment of whether an AI system is compatible with human autonomy, fundamental rights, transparency, accountability, non-discrimination, and effective human oversight. It may also serve as an interpretative criterion when existing legal rules are insufficiently specific to address new forms of technological mediation.

The construct therefore seeks to bridge the gap between the axiological dimension of human dignity and the operational requirements of AI governance. Its regulatory significance lies in requiring that technological innovation remain subordinated to the legal status of the person as a holder of rights and not merely as a source of data, an object of classification, or the passive recipient of automated decisions.

The principal legal challenge identified by the study is therefore not to determine whether artificial intelligence should be accepted or rejected as a technology. Rather, it is to establish the legal, ethical, and institutional conditions under which AI systems may be developed and used without undermining human dignity or the effective protection of fundamental rights.

The findings ultimately support a model of AI governance based on prevention, regulatory adaptability, multilevel cooperation, meaningful human oversight, transparency, accountability, non-discrimination, and effective mechanisms of review. Within this model, systemic dignity operates as the normative point of articulation between technological development and the continued primacy of the human person within the legal order.

Conclusions

The study demonstrates that the interaction between artificial intelligence, human dignity, and fundamental rights generates legal and technological challenges that increasingly exceed the response capacity of traditional regulatory frameworks. The rapid evolution of AI systems and their progressive incorporation into legally sensitive areas require legal responses capable of combining technological innovation with effective safeguards for the person as a holder of fundamental rights.

Human dignity remains the normative foundation of this relationship. Its protection cannot be confined to an abstract axiological declaration, but must inform the design, development, deployment, use, supervision, and review of artificial intelligence systems, particularly where automated decision-making may affect autonomy, privacy, equality, non-discrimination, access to opportunities, due process, and other legally protected interests.

The findings confirm that meaningful human oversight constitutes an essential legal safeguard. Automated systems should not displace human judgment in contexts involving significant legal consequences, nor should technological efficiency prevail over fundamental rights. Transparency, accountability, contestability, bias prevention, and access to effective review mechanisms therefore emerge as indispensable elements of rights-based AI governance.

The analysis also shows that the regulatory challenge is not limited to the absence of legislation. It also concerns the capacity of legal systems to translate general principles into concrete obligations, institutional responsibilities,

compliance mechanisms, and enforceable safeguards. In this respect, the risk-based approach adopted by Regulation (EU) 2024/1689 represents an important development because it links technological risk to the protection of health, safety, and fundamental rights and strengthens the juridical relevance of human oversight.

The paradigm of complexity confirms that artificial intelligence, law, ethics, technology, and society cannot be treated as isolated domains. Their interaction generates a sociotechnical environment in which legal norms influence technological design while technological developments simultaneously create new demands for legal interpretation, regulation, and institutional control.

Within this framework, the principal theoretical contribution of the study is the construct of this normative category. Systemic dignity is understood as the dynamic projection of human dignity into complex sociotechnical environments in which individuals, legal systems, and artificial intelligence technologies interact. It does not replace the classical human-centered conception of dignity, but extends its explanatory and regulatory function to contexts increasingly mediated by automated systems.

This framework may therefore operate both as a foundational principle and as an interpretative and regulatory criterion. Its function is to ensure that technological innovation remains subordinated to human autonomy, fundamental rights, legal accountability, transparency, non-discrimination, and effective human control throughout the AI life cycle.

The study ultimately concludes that the principal legal challenge is not to determine whether artificial intelligence is inherently beneficial or harmful, but to establish the normative, institutional, and ethical conditions under which its use may be considered legitimate. A human-rights-based model of AI governance requires preventive regulation, multilevel cooperation, adaptive legal frameworks, meaningful human oversight, enforceable accountability, and effective mechanisms for reviewing automated decisions.

Under this approach, this normative category provides a conceptual basis for preserving the continued primacy of the human person within the legal order while allowing technological development to advance under conditions compatible with fundamental-rights protection.

References

- Akula, R., & Garibay, I. (2021). Ethical AI for social good. In C. Stephanidis et al. (Eds.), *HCI International 2021—Late breaking papers: Multimodality, eXtended reality, and artificial intelligence* (Lecture Notes in Computer Science, Vol. 13095). Springer.
https://doi.org/10.1007/978-3-030-90963-5_28

- Arrestegui, L. B. (2012). Fundamentos históricos y filosóficos de la inteligencia artificial. UCV-HACER. *Revista de Investigación y Cultura*, 1(1), 87–92.
- Bach, T. A., Kaarstad, M., Solberg, E., et al. (2025). Insights into suggested Responsible AI (RAI) practices in real-world settings: A systematic literature review. *AI and Ethics*, 5, 3185–3232. <https://doi.org/10.1007/s43681-024-00648-7>
- Belouali, S., Belouali, A., Saber, M., Jaafar, K., & Belkasmi, M. G. (2020). Ethics of AI or ethical AI, topical point of view. In A. El Moussati, K. Kpalma, M. G. Belkasmi, M. Saber, & S. Guégan (Eds.), *Advances in smart technologies applications and case studies (Lecture Notes in Electrical Engineering, Vol. 684)*. Springer. https://doi.org/10.1007/978-3-030-53187-4_58
- Bostrom, N., & Yudkowsky, E. (2018). The ethics of artificial intelligence. In R. V. Yampolskiy (Ed.), *Artificial intelligence safety and security* (pp. 57–69). Chapman and Hall/CRC.
- Brey, P., & Dainow, B. (2024). Ethics by design for artificial intelligence. *AI and Ethics*, 4, 1265–1277. <https://doi.org/10.1007/s43681-023-00330-4>
- Brundage, M. (2015). Taking superintelligence seriously: Superintelligence: Paths, dangers, strategies by Nick Bostrom. *Futures*, 72, 32–35.
- Bryson, J. J. (2016). Patience is not a virtue: AI and the design of ethical systems. In *AAAI Spring Symposium Series*. Association for the Advancement of Artificial Intelligence.
- Cáceres Nieto, E. (2007). La teoría sintáctica en el naciente constructivismo jurídico. En *Constructivismo jurídico y metateoría del derecho*. Universidad Nacional Autónoma de México. <http://ru.juridicas.unam.mx/xmlui/handle/123456789/1431>
- Coeckelbergh, M. (2020). *AI ethics*. MIT Press.
- Consejo General del Poder Judicial. (2024). *Inteligencia artificial y justicia (Formación a Distancia, 2)*. Centro de Documentación Judicial.
- Dignum, V. (2019). *Responsible artificial intelligence: How to develop and use AI in a responsible way*. Springer.
- Dworkin, R. (1986). *Law's empire*. Harvard University Press.
- Esquivel García, R., Vargas Mursulí, F. M., & Navarrete Pita, Y. (2026). *Gobernanza algorítmica y People Analytics: un modelo para proteger los derechos laborales en organizaciones digitales*. Estudios Del Desarrollo Social: Cuba Y América Latina, 14, 1330–1347. <https://revistas.uh.cu/revflasco/article/view/12850>
- Eubanks, V. (2018). *Automating inequality: How high-tech tools profile, police, and punish the poor*. St. Martin's Press.
- European Parliament, & Council of the European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- Floridi, L. (2013). *The ethics of information*. Oxford University Press.
- Floridi, L., & Cows, J. (2021). A unified framework of five principles for AI in society. In L. Floridi (Ed.), *Ethics, governance, and policies in artificial intelligence* (pp. 5–17). Springer. https://doi.org/10.1007/978-3-030-81907-1_2
- Floridi, L., Cows, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People—An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28, 689–707. <https://doi.org/10.1007/s11023-018-9482-5>
- Habermas, J. (2010). El concepto de dignidad humana y la utopía realista de los derechos humanos. *Diánoia*, 55(64), 3–25.
- Hagendorff, T. (2022). Blind spots in AI ethics. *AI and Ethics*, 2(4), 851–867. <https://doi.org/10.1007/s43681-021-00122-8>
- Heilinger, J.-C. (2022). The ethics of AI ethics: A constructive critique. *Philosophy & Technology*, 35, Article 61. <https://doi.org/10.1007/s13347-022-00557-9>
- Hernández, F. M. Á., Picarella, L., & Fiorino, V. R. M. (2024). Los desafíos éticos-jurídicos de la inteligencia artificial en Europa y Colombia. *Clío. Revista de Historia, Ciencias Humanas y Pensamiento Crítico*, (9), 866–907.
- Kant, I. (2012). *Fundamentación para una metafísica de las costumbres* (2.ª ed.). Alianza Editorial.
- Llano Alonso, F. H., Garrido Martín, J., & Valdivia Jiménez, R. (Coords.). (2022). *Inteligencia artificial y filosofía del derecho*. Ediciones Laborum.
- Morin, E. (2002). *La cabeza bien puesta: Repensar la reforma, reformar el pensamiento* (1.ª ed., 5.ª reimp.). Nueva Visión.
- Müller, V. C. (2023). Ethics of artificial intelligence and robotics. In E. N. Zalta & U. Nodelman (Eds.), *The Stanford encyclopedia of philosophy* (Fall 2023 ed.). Stanford University. <https://plato.stanford.edu/archives/fall2023/entries/ethics-ai/>
- Orwat, C., Bareis, J., Folberth, A., Jähnel, J., & Wadephul, C. (2024). Normative challenges of risk regulation of artificial intelligence. *NanoEthics*, 18, Article 11. <https://doi.org/10.1007/s11569-024-00454-9>
- Russell, S. (2019). *Human compatible: Artificial intelligence and the problem of control*. Viking.
- Russell, S. J., & Norvig, P. (2016). *Artificial intelligence: A modern approach* (3rd ed.). Pearson.
- Ruttkamp-Bloem, E. (2020). The quest for actionable AI ethics. In A. Gerber (Ed.), *Artificial intelligence research: First Southern African Conference for AI*

- Research, SACAIR 2020 (pp. 34–50). Springer.
https://doi.org/10.1007/978-3-030-66151-9_3
- Shaelou, S. L., & Razmetaeva, Y. (2023). Challenges to fundamental human rights in the age of artificial intelligence systems: Shaping the digital legal order while upholding rule of law principles and European values. *ERA Forum*, 24, 567–587.
<https://doi.org/10.1007/s12027-023-00777-2>
- Sun, W. (2018). Artificial intelligence and ethical principles. In D. Jin (Ed.), *Reconstructing our orders*. Springer. https://doi.org/10.1007/978-981-13-2209-9_2
- Teo, S. A. (2025). Artificial intelligence and its ‘slow violence’ to human rights. *AI and Ethics*, 5, 2265–2280. <https://doi.org/10.1007/s43681-024-00547-x>
- Thielscher, C. (2025). Dignity as a concept for computer ethics. *AI and Ethics*, 5, 4061–4067.
<https://doi.org/10.1007/s43681-025-00693-w>
- Turing, A. M. (2009). Computing machinery and intelligence. In R. Epstein, G. Roberts, & G. Beber (Eds.), *Parsing the Turing test: Philosophical and methodological issues in the quest for the thinking computer* (pp. 23–65). Springer.
- UNESCO. (2021). Recommendation on the Ethics of Artificial Intelligence.
<https://www.unesco.org/es/legal-affairs/recommendation-ethics-artificial-intelligence>
- United Nations. (1948). Universal Declaration of Human Rights. <https://www.un.org/es/about-us/universal-declaration-of-human-rights>
- United Nations. (1966a). International Covenant on Civil and Political Rights.
<https://www.ohchr.org/es/instruments-mechanisms/instruments/international-covenant-civil-and-political-rights>
- United Nations. (1966b). International Covenant on Economic, Social and Cultural Rights.
<https://www.ohchr.org/es/instruments-mechanisms/instruments/international-covenant-economic-social-and-cultural-rights>
- Wong, A. (2021). Ethics and regulation of artificial intelligence. In E. Mercier-Laurent, M. Ö. Kayalica, & M. L. Owoc (Eds.), *Artificial intelligence for knowledge management: AI4KM 2021 (IFIP Advances in Information and Communication Technology, Vol. 614)*. Springer.
https://doi.org/10.1007/978-3-030-80847-1_1
- Yeung, K. (2018). Algorithmic regulation: A critical interrogation. *Regulation & Governance*, 12(4), 505–523. <https://doi.org/10.1111/regg.12158>

Conflicts of interest

The author declares that he has no conflicts of interest.

Author contributions

Javier A. Artilles-Santana: Conceptualization, data curation, formal analysis, investigation, methodology, supervision,

validation, visualization, drafting the original manuscript and writing, review, and editing.

Data availability statement

The datasets used and/or analyzed during the current study are available from the corresponding author on reasonable request.

Statement on the use of AI

The author acknowledges the use of generative AI and AI-assisted technologies to improve the readability and clarity of the article.

Disclaimer/Editor’s note

The statements, opinions, and data contained in all publications are solely those of the individual authors and contributors and not of Journal of Law and Epistemic Studies.

Journal of Law and Epistemic Studies and/or the editors disclaim any responsibility for any injury to people or property resulting from any ideas, methods, instructions, or products mentioned in the content.