

HUMAN DIGNITY AND ARTIFICIAL INTELLIGENCE: LEGISLATIVE PATHWAYS FOR PROTECTING DIGNITY IN THE AI ERA

Lakshay*
Prof. Devinder Singh**

ABSTRACT

Artificial intelligence, a rapidly evolving technology, is transforming multiple industries by improving efficiency, streamlining tasks and fostering the development of innovative products, yet this accelerated growth raises considerable ethical concern regarding the impact of these technologies on human dignity. This paper examines the opportunities and challenges inherent in the evolving fields of artificial intelligence and robotics in relation to human dignity, situating the concept within its constitutional and international legal history before assessing how emerging technology places pressure upon it. Through critical analysis of existing literature, statutory instruments and comparative regulatory practice, the paper examines the ethical considerations most consequential for human dignity, algorithmic bias, autonomy and control, and privacy, and surveys the legislative and regulatory frameworks that have emerged in India and internationally in response. The paper argues that while the advancement of artificial intelligence is unlikely to be halted and has already produced considerable benefit across sectors, existing Indian law remains substantially reliant on general-purpose statutes and non-binding guidance rather than a dedicated framework calibrated specifically to dignity-related harm, and it proposes concrete legislative and institutional pathways capable of closing that gap.

KEYWORDS: *Artificial Intelligence; Human Dignity; Human Rights; Ethical AI; AI Regulation*

1. INTRODUCTION

The concept of dignity, and the diverse principles it encompasses, forms a foundational element of international human rights discourse and of the ethical systems that protect individuals from one another and from the exercise of state power. Dignity may be invoked implicitly, as in the United States Declaration of Independence's assertion that all men are created equal, or explicitly, as in the Universal Declaration of Human Rights' recognition that the inherent dignity and the equal and inalienable rights of all members of the human family form the foundation of freedom, justice and peace in the world.¹ Dignity functions as a foundational element not only for the protection of the individual but, as the opening lines of the United Nations Charter itself suggest, for the rules-based international order established after the Second World War; the 1919 Covenant of the League of Nations, by contrast, sought to establish peace and security without taking dignity or individual rights into account, and the comparative absence of a rights-anchored international order in the interwar

* Research Scholar, Department of Laws, Panjab University, Chandigarh.

** Professor & Vice Chancellor, Dr. B.R. Ambedkar National Law University, Sonapat.

¹ Universal Declaration of Human Rights pmbl., G.A. Res. 217A (III) (Dec. 10, 1948).

period supplies a cautionary historical benchmark against which the post-1945 dignity-centred order should be measured.²

This historical trajectory matters directly for the analysis that follows, since it establishes dignity not as a recently invented rhetorical device but as a foundational commitment whose absence, in the interwar period, coincided with catastrophic institutional failure. The question this paper examines is whether a comparable gap is now opening between the pace of technological change and the institutional capacity of law to embed dignity within it, a gap analogous in structure, though obviously not in scale or consequence, to the gap between the League of Nations' security-focused mandate and the rights-based order that followed it. Artificial intelligence, on this reading, represents merely the latest and most consequential test of whether legal institutions can keep the dignity commitment operative as the technological and institutional context around it changes.

As attention to the effects of artificial intelligence on human rights has grown, human dignity has emerged as a fundamental value informing the governance of new and emerging technology, reflecting a broader concern that existing legal frameworks may not adequately address the consequences such technology can produce, and that a return to foundational principles, including respect for human dignity, is accordingly necessary.³ Artificial intelligence, as originally defined by John McCarthy, who coined the term in 1956, denotes the science and engineering of making intelligent machines, the simulation of human intelligence processes by computer systems capable of reasoning, discovering meaning or learning from experience.⁴ Legislators, policymakers and civil society organisations have increasingly called for measures addressing the threat such systems may pose to human dignity specifically, situating dignity alongside individual freedom, democracy, justice, the rule of law, equality, non-discrimination and solidarity as a foundational principle any AI risk assessment must engage. While the rapid advancement of new technology poses genuine challenges to the practical realisation of dignity, its fundamental meaning remains unchanged: every individual possesses an intrinsic value that must never be endangered or suppressed, and any action inconsistent with that value constitutes a violation regardless of the technological means through which it occurs.

The notion of human dignity gained prominence in international law through its appearance in the Preamble to the United Nations Charter, which

² Jon Rueda et al., *Why Dignity Is a Troubling Concept for AI Ethics*, 6(3) *Patterns* 101207 (2025).

³ Sheila Jasanoff, *The Ethics of Invention: Technology and the Human Future* (1st ed., W.W. Norton & Co. 2016).

⁴ John McCarthy, *What Is Artificial Intelligence?* (Stanford Univ. 2007).

sought to reaffirm faith in fundamental human rights, in the dignity and worth of the human person, and in the equal rights of men and women and of nations large and small.⁵ The Charter and the Universal Declaration of Human Rights together emphasised the concept, and numerous national constitutions have since recognised it; although dignity remains formally undefined in international instruments, it is understood as a fundamental value inherent to all human beings from which corresponding rights and obligations flow. Concern regarding the risks artificial intelligence poses to this value is not confined to critics of the technology; it has been voiced by figures central to the field's own development, including Geoffrey Hinton, recipient of the 2024 Nobel Prize in Physics for his foundational work on deep learning and neural networks, whose research supplied the intellectual foundation for the generative AI systems major technology companies now regard as central to their future, and who has himself warned publicly of the risks that same technology poses.

1.1 Objectives of the Study

This paper critically examines how existing and emerging legislative frameworks can safeguard human dignity in the context of artificial intelligence, ensuring that technological innovation aligns with fundamental rights and ethical principle, and explores potential pathways for developing inclusive, forward-looking law capable of balancing artificial intelligence's societal benefit against the need to protect individuals from dignity-related risk, including bias, surveillance and the erosion of personal autonomy.

1.2 Research Methodology

The study adopts a doctrinal approach supported by an analytical research framework, relying on qualitative methods to systematically examine how artificial intelligence bears upon the protection of human dignity within legislative and regulatory contexts, drawing on Indian statutory and constitutional material alongside comparative international instruments.

2. THE CONCEPT OF HUMAN DIGNITY: DOCTRINAL FOUNDATIONS

Considerable debate persists over precisely which conception of dignity is invoked when artificial intelligence is said to threaten it. On one account, dignity is understood as an intrinsic, inviolable quality of human beings, unconditional and incapable of diminishment; on this understanding, it is difficult to see how a technology could threaten dignity so conceived, since intrinsic worth cannot, by definition, be taken away, though it may certainly be disregarded or overlooked by those who treat a person as though they lacked it. A second, relational conception treats dignity as a form of social status capable

⁵ U.N. Charter pmbl.

of degradation and humiliation, since it depends on characteristics or treatment not inherent to the person as such; it is on this second, status-based conception that claims about artificial intelligence's threat to dignity become considerably more plausible, since specific, identifiable harms, humiliating automated decision-making, discriminatory profiling, degrading surveillance, can meaningfully be assessed against a relational standard in a way they cannot against an unconditional one.⁶

This conceptual ambiguity carries a further, more practical implication worth making explicit: appeals to dignity in AI ethics discourse frequently function as a proxy for other, more specific ethical claims, equality, non-discrimination, autonomy, or the prevention of undue harm, rather than as an independent standard in their own right.⁷ Recognising this proxy function is analytically useful rather than reductive, since it allows the discussion that follows to identify, for each specific technological harm considered, which underlying value, equality, autonomy, privacy, dignity-talk is actually protecting, and to assess the adequacy of existing law against that more precise standard rather than against the admittedly elastic language of dignity alone.⁸

3. ETHICAL DIMENSIONS OF ARTIFICIAL INTELLIGENCE AND THEIR IMPLICATIONS FOR DIGNITY

Ensuring ethical standards in the development and use of artificial intelligence requires the establishment of responsible AI principles, an endeavour to which initiatives including the IEEE Global Initiative on Ethics of Autonomous and Intelligent Systems and the AI4People project have contributed substantially, centring AI ethics on fairness, transparency, accountability and privacy as principles requiring integration into the design, development and deployment of AI systems rather than treatment as after-the-fact compliance concerns. UNESCO's Recommendation on the Ethics of Artificial Intelligence identifies particular attention as warranted across several domains: education, where digital societies require innovative pedagogical practice, ethical reflection and skills aligned with labour-market demand; science, where AI introduces new research capability while reshaping our understanding of scientific explanation itself; cultural identity and diversity, where AI technology may enhance cultural and creative industry even as it risks concentrating cultural content, data and income among a limited number of actors, with adverse consequences for linguistic and cultural pluralism; and communication and information, where the algorithmic architecture governing content moderation and curation on social media and search platforms raises acute concern regarding access to

⁶ L. Zardiashvili & E. Fosch-Villaronga, "Oh, Dignity Too?" Said the Robot: Human Dignity as the Basis for the Governance of Robotics, 22 *Minds & Machines* 121 (2020).

⁷ R. Macklin, Dignity Is a Useless Concept, 327 *Brit. Med. J.* 1419 (2003).

⁸ *Id.*

information, disinformation, discrimination, freedom of expression and media literacy.⁹

Algorithmic bias represents perhaps the most extensively documented dignity-related harm associated with contemporary AI systems: despite diligent effort, algorithms can reinforce and amplify prejudice already present within their training data, producing discriminatory outcomes against underrepresented groups, a pattern most visibly documented in facial recognition software, which has repeatedly demonstrated lower accuracy for people of colour, with consequent risk of wrongful identification and downstream legal or reputational harm.¹⁰ The dignity implication of this pattern extends beyond the immediate risk of misidentification: where a technology systematically performs less reliably for a particular demographic group, its deployment in consequential contexts, policing, hiring, credit assessment, effectively subjects that group to a materially different, and inferior, standard of automated decision-making than the population for whom the system was principally trained and validated, reproducing through technical means precisely the kind of unequal treatment the equality guarantees examined in Part 5 are designed to prevent.

A second challenge concerns the balance between autonomy and control: AI systems making decisions independent of human oversight raise genuine questions of accountability, exemplified by the difficulty of assigning responsibility where an autonomous vehicle's decision produces harm, a difficulty underscoring the more general principle that AI systems should be designed to augment human decision-making rather than to supplant it entirely.¹¹ This autonomy concern carries a dignity dimension distinct from, though related to, the accountability question: a person subject to a fully automated decision with no meaningful opportunity for human review is, in a real sense, denied the individualised consideration that dignity-respecting treatment ordinarily requires, since the automated system, however sophisticated, applies a general rule to the individual case without the capacity for the contextual judgment a human decision-maker could in principle bring to bear.

A third concern addresses privacy: AI's capacity to analyse extensive personal data raises genuine questions of surveillance and consent, illustrated by smart home devices capable of inadvertently recording private conversation, such that upholding individual privacy against this capability becomes a precondition for sustaining public trust in the technology more broadly.¹² The scale at which contemporary AI systems operate compounds this privacy

⁹ UNESCO, Recommendations on the Ethics of Artificial Intelligence (2021).

¹⁰ Maryama Elmi, Faces and Places: Navigating the Cross-Border Legal Challenges of AI Facial Recognition Technologies, 5 AI & Ethics 5453 (2025).

¹¹ Id.

¹² Id.

concern considerably beyond what earlier, narrower data-processing technology presented: a single AI system may draw on data aggregated across millions of individuals to train a model subsequently applied to any one of them, meaning that a person's exposure to algorithmic decision-making is shaped not only by data they themselves provided but by the aggregate patterns the system extracted from a population to which they may have contributed only indirectly or not at all, a structural feature of machine learning that existing privacy law, built around the more individualised paradigm of direct data collection and use, does not straightforwardly address. Designing AI systems that align with these core values requires embedding ethical decision-making capacity directly into system architecture, promoting trust through transparency and accountability rather than treating these properties as attributes to be added once a system's core function has already been finalised.¹³

4. INDIA'S STATUTORY AND REGULATORY FRAMEWORK

India's foundational digital-governance statute, the Information Technology Act, 2000, provides legal recognition to electronic records and digital signatures and has evolved over time to address emerging challenges including those artificial intelligence and data-driven technology present.¹⁴ Section 43A places liability on companies handling sensitive personal data to implement reasonable security practice, a provision of particular relevance to AI systems processing large volumes of personal information including biometric or financial data, while Section 66 addresses computer-related offences including hacking and unauthorised access, offences of particular significance where automated systems are exploited for malicious purpose.¹⁵ These provisions, drafted well before contemporary machine-learning technology existed in its present form, apply to AI-related harm only by extension rather than by specific design, a limitation that recurs throughout India's present statutory landscape and that Part 7 returns to directly. The National Strategy for Artificial Intelligence, released by NITI Aayog in 2018, remains India's first comprehensive AI policy framework; though a non-binding policy document rather than law, it identifies agriculture, healthcare, education, smart cities and mobility as priority sectors, emphasises AI research, data-sharing infrastructure and skilled human capital, and treats ethical deployment as a core organising principle intended to keep the resulting AI ecosystem inclusive, transparent and aligned with India's democratic values.¹⁶

The Information Technology (Intermediary Guidelines and Digital Media Ethics Code) Rules, 2021, issued under the parent Act, impose due-diligence

¹³ Id.

¹⁴ The Information Technology Act, No. 21 of 2000 (India).

¹⁵ The Information Technology Act, No. 21 of 2000, §§ 43A, 66 (India).

¹⁶ NITI Aayog, National Strategy for Artificial Intelligence: #AIforAll (June 2018).

obligations on online intermediaries including social media platforms and AI-driven services, mandating grievance-redressal mechanisms, compliance officers and prompt removal of flagged content within specified timelines, and extending a three-tier regulatory structure to digital media and OTT platforms specifically.¹⁷ The Digital Personal Data Protection Act, 2023 establishes rights for data principals and obligations for data fiduciaries, requiring that AI systems processing personal data do so lawfully, transparently and for a specified purpose, while introducing cross-border transfer rules and penalties for misuse of particular relevance to AI deployment in healthcare, finance and e-commerce. The Act's data-fiduciary obligations apply generally to any entity processing personal data rather than specifically to AI systems, meaning that the distinctive risks machine-learning systems present, inferential profiling from lawfully collected data, opaque automated decision-making, algorithmic discrimination through proxy variables, remain addressed only indirectly, through the general data-protection architecture, rather than through provisions calibrated to the specific technical character of AI processing. Most recently, the India AI Governance Guidelines, released by the Ministry of Electronics and Information Technology in 2025, articulate seven foundational principles, fairness, reliability, explainability, ethics, accountability, inclusivity and sustainability, functioning as a soft-law instrument offering direction and best practice pending the enactment of comprehensive, binding AI-specific legislation.¹⁸ The Guidelines' explicit acknowledgment that binding legislation remains a future rather than present undertaking confirms the central structural finding this paper advances: India presently governs artificial intelligence's dignity-related risk almost entirely through a combination of general-purpose statute and non-binding guidance, a combination this paper's comparative analysis in Part 6 shows falls short of the binding, AI-specific regulatory models several other major jurisdictions have already adopted.

5. CONSTITUTIONAL FOUNDATIONS OF HUMAN DIGNITY IN INDIA

India's constitutional framework reflects its obligations under international human rights law by safeguarding personal autonomy and dignity through several interlocking provisions. Article 21 guarantees the right to life and personal liberty, protecting individuals against arbitrary state action, and has been expansively interpreted by the judiciary to encompass privacy, autonomy and dignity as integral components of the right to life itself.¹⁹ Article 14 reinforces the principle of equality before the law, prohibiting discrimination

¹⁷ Information Technology (Intermediary Guidelines and Digital Media Ethics Code) Rules, 2021 (India).

¹⁸ Ministry of Electronics & Information Technology, India AI Governance Guidelines (2025).

¹⁹ The Constitution of India, 1950, art. 21.

and mandating equal protection in a manner that prevents practices undermining individual dignity and ensures fairness in governance more broadly.²⁰ Article 19 safeguards freedom of speech and expression, empowering individuals to assert identity, belief and dignity without arbitrary restriction, a right of particular contemporary significance where digital governance and AI-mediated content moderation intersect with expressive freedom.²¹ Together, these provisions establish a constitutional framework within which the legal discourse traced in this paper concerning artificial intelligence and digital governance must necessarily operate. The Protection of Human Rights Act, 1993 further strengthens this constitutional commitment by defining human rights as those relating to life, liberty, equality and dignity and by establishing the National Human Rights Commission, a statutory body empowered to investigate rights violations, propose remedial measures and monitor compliance, addressing concerns including custodial violence, discrimination and abuse of authority in a manner that reinforces the constitutional promise of dignity independent of, and in addition to, the specific AI-related statutes examined in Part 4.²²

6. INTERNATIONAL AND COMPARATIVE REGULATORY FRAMEWORKS

Internationally, several institutions have advanced comprehensive frameworks for AI governance with direct implications for dignity. The European Union's Artificial Intelligence Act, Regulation (EU) 2024/1689, represents the world's first comprehensive, binding AI statute, and it protects human dignity directly by prohibiting unacceptable-risk practices including systems that deploy subliminal or manipulative techniques impairing free will and systems that exploit the vulnerability of specific groups, defined by age, disability or comparable characteristic, to distort behaviour.²³ The Act's risk-tiered structure, extending graduated obligation to high-risk systems short of outright prohibition, offers a template India's own regulatory design could adapt, since it demonstrates that dignity-protective regulation need not proceed through blanket prohibition alone but can instead calibrate obligation to the severity of risk a given application presents, preserving space for lower-risk AI application to develop with comparatively lighter regulatory burden.

The Council of Europe's Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, CETS No. 225, adopted by the Committee of Ministers on 17 May 2024 and opened for signature on 5

²⁰ The Constitution of India, 1950, art. 14.

²¹ The Constitution of India, 1950, art. 19.

²² The Protection of Human Rights Act, No. 10 of 1994, § 3 (India).

²³ Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 (Artificial Intelligence Act), 2024 O.J. (L 1689).

September 2024, constitutes the first legally binding international treaty in the field, requiring signatories to ensure that AI systems uphold human rights, democracy and the rule of law, and expressly mandating attention to human dignity and individual autonomy across the entire AI lifecycle.²⁴ The Convention's deliberately open character, extending signature eligibility to non-member states and international organisations beyond the Council of Europe's own membership, reflects a considered ambition toward genuinely global standard-setting rather than a regional instrument alone, and its ratification to date by a comparatively modest number of states illustrates both the significance of achieving binding international consensus on AI governance and the practical difficulty of securing it. The European Union's General Data Protection Regulation contributes a further dignity-protective mechanism through Article 22, which grants individuals the right to object to automated decision-making that produces legal or similarly significant effects, ensuring that a person is not simply reduced to an algorithmic outcome without avenue for human review.²⁵

Beyond the European context, UNESCO's Recommendation on the Ethics of Artificial Intelligence, while not legally binding, applies across the organisation's 194 member states and centres AI development on the protection of human rights and dignity, with particular emphasis on sustained human oversight throughout the AI lifecycle.²⁶ The Recommendation's near-universal membership base gives it a breadth of formal endorsement no binding treaty presently matches, even though that breadth is achieved precisely by forgoing binding legal obligation, illustrating a recurring trade-off in international AI governance between the depth of commitment a binding instrument secures and the breadth of adoption a non-binding one can achieve. The United States' Blueprint for an AI Bill of Rights, released by the White House Office of Science and Technology Policy in 2022, outlines measures directed at ensuring the safety and effectiveness of AI systems and at safeguarding human autonomy, functioning, as with India's own 2025 Governance Guidelines, as a non-binding statement of principle rather than enforceable law.²⁷ Read together, these instruments demonstrate a common international convergence on the substantive principles, transparency, fairness, human oversight and accountability, that dignity-protective AI regulation requires, even as jurisdictions diverge considerably in whether and how they have translated those principles into binding legal obligation, with the European Union's twin

²⁴ Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, CETS No. 225 (opened for signature Sept. 5, 2024).

²⁵ Regulation (EU) 2016/679 (General Data Protection Regulation), art. 22, 2016 O.J. (L 119) 1.

²⁶ UNESCO, Recommendation on the Ethics of Artificial Intelligence (Nov. 24, 2021).

²⁷ White House Office of Science & Technology Policy, Blueprint for an AI Bill of Rights (Oct. 2022).

instruments, the AI Act and the Council of Europe Framework Convention, currently representing the most advanced binding models against which India's own predominantly soft-law approach, traced in Part 4, may usefully be measured.

7. TOWARDS A DIGNITY-CENTRED LEGISLATIVE FRAMEWORK: RECOMMENDATIONS

The comparative analysis above points toward several concrete measures capable of strengthening India's dignity-protective AI governance. Parliament should enact binding regulation establishing clear obligations for AI development and deployment, moving beyond the presently voluntary 2025 Governance Guidelines toward the kind of enforceable framework the European Union's AI Act supplies, while embedding mandatory human oversight for high-stakes automated decisions in domains including healthcare, criminal justice and employment, ensuring individuals are not reduced to algorithmic outcomes without meaningful recourse to human review, following the model Article 22 of the GDPR already supplies in the European context. A risk-tiered structure modelled on, but not identical to, the European Union's own approach would allow such legislation to calibrate obligation to the severity of risk a given application presents, avoiding the twin failure modes of either uniform light-touch regulation that leaves genuinely high-risk applications inadequately governed or uniform heavy-touch regulation that unduly burdens lower-risk innovation the Indian digital economy continues to depend upon.

Independent ethics review bodies should be established to oversee significant AI projects, particularly within government, ensuring compliance with dignity-centred standards through a mechanism structurally independent of the developing agency itself, while regular bias auditing and fairness testing of datasets and algorithms should be mandated to detect and mitigate the discriminatory outcomes traced in Part 3, particularly within facial recognition and other biometric applications where the documented accuracy disparities carry the most direct bearing on individual dignity and liberty. Such audit requirements should extend specifically to public-sector deployment, including law enforcement and welfare administration, where the consequences of undetected bias fall on citizens with the least practical capacity to contest an adverse automated determination.

A statutory right to explanation, guaranteeing individuals a meaningful account of how an AI system reached a decision materially affecting them, would reinforce both autonomy and accountability in a manner presently absent from Indian law outside the DPDP Act's general transparency obligations, and should be paired with accessible redress mechanisms through which individuals may challenge or seek remedy against harmful automated decisions. Data minimisation and privacy protection should be strengthened specifically in

relation to AI training and deployment, limiting collection to what use genuinely requires rather than what convenience or future flexibility might suggest, while ethics and human-rights training should be integrated into AI professional education to cultivate a durable culture of responsibility among developers themselves rather than relying exclusively on external regulatory enforcement. Finally, India should pursue closer engagement with the international instruments surveyed in Part 6, including consideration of eventual accession to the Council of Europe's Framework Convention, since global cooperation on harmonised AI standards would help ensure that dignity protections remain consistent across the borders AI systems and the data they process routinely cross, rather than varying unpredictably according to the jurisdiction in which a given system happens to be deployed.

8. CONCLUSION

Respecting human dignity in AI ethics is essential to ensuring that AI systems promote human well-being and respect human autonomy, and prioritising transparency, accountability and human-centred design offers the most direct route to mitigating the risks this paper has traced while preserving the technology's genuine benefit. The advancement of AI systems that uphold human dignity requires a genuinely multidisciplinary approach engaging technologists, ethicists, lawyers and social scientists together, since no single discipline can adequately address the full range of concern artificial intelligence raises; legal experts in particular occupy a distinctive position in ensuring that AI is developed and deployed in a manner consistent with human rights, and governments, as the primary developers and regulators of this technology within their own jurisdictions, bear a corresponding responsibility to ensure its application remains just and equitable. Regular evaluation and audit of AI systems, together with the development of dignity-specific human rights impact assessments capable of anticipating a system's effects before deployment rather than merely responding to harm after the fact, would meaningfully advance this objective. India's existing framework, comprising the Information Technology Act, the Digital Personal Data Protection Act, the Intermediary Guidelines and the 2025 Governance Guidelines, together with the constitutional foundation Articles 14, 19 and 21 supply, already contains the doctrinal resources from which a genuinely binding, dignity-centred AI statute could be constructed; what remains outstanding, as this paper has argued throughout, is the legislative act that would convert those resources into enforceable law before, rather than only after, artificial intelligence has caused the harm such a statute exists to prevent. It remains, in the final analysis, humanity's responsibility to govern this technology rather than to be governed by it.

BIBLIOGRAPHY

Books

1. Virginia Dignum, *Responsible Artificial Intelligence: How to Develop and Use AI in a Responsible Way* (Springer 2019).
2. Luciano Floridi, *The Ethics of Artificial Intelligence: Principles, Challenges, and Opportunities* (Oxford Univ. Press 2023).
3. Henry A. Kissinger, Eric Schmidt & Daniel Huttenlocher, *The Age of AI: And Our Human Future* (John Murray 2024).
4. Stuart Russell, *Human Compatible: AI and the Problem of Control* (Penguin 2020).
5. Max Tegmark, *Life 3.0: Being Human in the Age of Artificial Intelligence* (Knopf 2017).

Journal Articles

1. Jon Rueda et al., *Why Dignity Is a Troubling Concept for AI Ethics*, 6(3) *Patterns* 101207 (2025).
2. Nayef Al-Rodhan, *Artificial Intelligence: Implications for Human Dignity and Governance*, 27(3) *Oxford Pol. Rev.* 45 (2021).
3. Ginevra Le Moli, *AI vs. Human Dignity: When Human Underperformance Is Legally Required*, 4(1) *Revue Européenne du Droit* 105 (2022).
4. Sue Anne Teo, *Human Dignity and AI: Mapping the Contours and Utility of Human Dignity in Addressing Challenges Presented by AI*, 15(1) *L. Innovation & Tech.* 241 (2023).
5. Maryama Elmi, *Faces and Places: Navigating the Cross-Border Legal Challenges of AI Facial Recognition Technologies*, 5 *AI & Ethics* 5453 (2025).
6. R. Macklin, *Dignity Is a Useless Concept*, 327 *Brit. Med. J.* 1419 (2003).
7. L. Zardiashvili & E. Fosch-Villaronga, *"Oh, Dignity Too?" Said the Robot: Human Dignity as the Basis for the Governance of Robotics*, 22 *Minds & Machines* 121 (2020).

Statutes and Constitutional Materials

1. The Constitution of India, 1950, arts. 14, 19, 21.
2. The Information Technology Act, No. 21 of 2000, §§ 43A, 66 (India).
3. The Digital Personal Data Protection Act, No. 22 of 2023 (India).
4. The Protection of Human Rights Act, No. 10 of 1994, § 3 (India).
5. Information Technology (Intermediary Guidelines and Digital Media Ethics Code) Rules, 2021 (India).

International Instruments and Reports

1. U.N. Charter pmb.
2. Universal Declaration of Human Rights, G.A. Res. 217A (III) (Dec. 10, 1948).

3. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 (Artificial Intelligence Act), 2024 O.J. (L 1689).
4. Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, CETS No. 225 (opened for signature Sept. 5, 2024).
5. Regulation (EU) 2016/679 (General Data Protection Regulation), art. 22, 2016 O.J. (L 119) 1.
6. UNESCO, Recommendation on the Ethics of Artificial Intelligence (Nov. 24, 2021).
7. White House Office of Science & Technology Policy, Blueprint for an AI Bill of Rights (Oct. 2022).
8. NITI Aayog, National Strategy for Artificial Intelligence: #AIforAll (June 2018).
9. Ministry of Electronics and Information Technology, India AI Governance Guidelines (2025).