



Responsible Artificial Intelligence: A Comprehensive Review of Frameworks, Principles, and Practices

Shraddha Nitin Magdum¹ , *Tanuja Satish Dhope (Shendkar)² ,
Priyanka Kashinath Kendre³ , Gulnaz T. Thakur⁴ ,
Rajesh Kumar Kaushal⁵

¹Research Scholar (Electronics Engineering), Department of Electronics Engineering, Bharati Vidyapeeth (Deemed to be University) College of Engineering, Pune, Maharashtra, India.

²Department of Electronics & Communication Engineering, Bharati Vidyapeeth (Deemed to be University) College of Engineering, Pune, Maharashtra, India.

³Department of Information Technology, Trinity Academy of Engineering, Pune, Maharashtra, India.

⁴Department of Information Technology, Trinity Academy of Engineering, Pune, Maharashtra, India.

⁵Chitkara University Institute of Engineering and Technology, Chitkara University, Punjab, India

{shraddhagajil1993@gmail.com, *tanuja_dhope@yahoo.com,
priyankakendre.tae@kjei.edu.in, gulnazthakur.tae@kjei.edu.in,
rajesh.kaushal@chitkara.edu.in}

Abstract. The rapid proliferation of artificial intelligence systems across critical societal domains has intensified concerns about their ethical, social, and technical implications. This comprehensive review synthesizes recent scholarly literature on Responsible Artificial Intelligence (RAI), examining the multidimensional landscape of frameworks, principles, and practices that guide the development and deployment of trustworthy AI systems. Through systematic analysis of 30 recent publications (2024-2025), we identify and critically evaluate nine core dimensions of responsible AI: ethical frameworks and principles, fairness and bias mitigation, transparency and explain ability, accountability and governance, trustworthy AI systems, safety and robustness, privacy preservation, human-centered design, and regulatory approaches. Our analysis reveals significant progress in conceptual framework development and technical methods, yet identifies persistent challenges including the generalization gap in defenses against evolving threats, inadequate real-world evaluation protocols, fragmented regulatory landscapes, and the tension between theoretical principles and practical implementation. We synthesize emerging best practices, highlight critical research gaps, and propose an integrated research agenda that bridges technical rigor with social responsibility. This review contributes to the field by providing a holistic understanding of the responsible AI landscape and identifying pathways toward AI systems that are not only technically robust but also ethically aligned, socially beneficial, and publicly trustworthy.

Keywords: Responsible AI, AI Ethics, Fairness, Transparency, Accountability, Trustworthy AI, AI Governance, Explainability, AI Safety, Human-Centered AI

1 Introduction

Artificial intelligence (AI) has come a long way back into the niche of computer science research. It is now a potent and pervasive technology that is involved in major decision-making in numerous industries, such as in healthcare, criminal justice, banking and finance, employment and staffing, education, and even in autonomous vehicles and robotics [1]. Although the introduction of AI into these vital spheres has caused significant benefits, including the enhancement of efficiency, automation of large scale, and the ability to solve problems, it has also presented serious issues. These issues are associated with the issues of ethics, social impact, privacy, safety, and the safeguarding of the basic human rights. The more the AI systems will affect the lives of actual people, the more they can destroy the democratic processes and breed distrust should they be not developed and implemented in a responsible manner [2]. To solve these emerging issues, a concept of Responsible Artificial Intelligence (RAI) has appeared. RAI is a general and universal solution to make sure that AI technologies are created and implemented in a manner that is ethical, fair, transparent, accountable, safe, and focused on the human values and social expectations [3]. Responsible AI provides a comprehensive perspective unlike small technical approaches that only solve a single issue - e.g., accuracy improvement or bias reduction. It acknowledges that people put their faith in AI not only based on the level of performance of the system, but also on whether it adheres to societal norms, fairly treats people, keeps them safe, and can be comprehended and managed by humans [10].

Responsible AI has increasingly become a central concern of scholarly publications, political and governmental policy, and commercial practice over the past few years. This payoff of attention has been promoted in large part by a series of high profile instances of AI systems acting in unexpected manners, or favoring specific populations, abusing personal information, or making opaque and inexplicable decisions [4]. Consequently, various ethical codes of conduct, regulation policies, and technical standards have been designed by international organizations, governments, and technology corporations with a view to facilitate the responsible development of AI [5]. Although responsible AI practices are growing at a very fast pace, there are still critical issues. Most of the programs are still disjointed, employ high-level principles that can hardly be implemented, and base on poor evaluation procedures that cannot represent real world risks. There is also still a significant disparity between the purported practice and reality in AI development and deployment [6].

This review mentions these concerns by summarizing the recent 2024/2025 studies to provide a full perspective of the responsible AI situation. It frames the field in terms of nine dimensions that describe responsible AI: moral grounding in terms of ethical frameworks, mitigation of bias and fairness, transparency and explainability, governance and accountability, trustworthy AI, which implies reliability and value alignment, safety and resilience, preservation of privacy, human-centered design, and regu-

latory and policy frameworks [7]. Figure 1 shows multidimensional framework of responsible AI.

The review is valuable as it provides: (1) a synthesized summary of all of the major approaches and findings to date, (2) a critical assessment of existing strengths, weaknesses, and underdeveloped areas, (3) a sense of the gaps in research, as well as an integrated agenda in regard to the technical, ethical, and societal perspectives, and (4) practical recommendations that can help policymakers, researchers and industry professionals adopt the principles of responsible AI successfully. The rest of this paper will be organized as follows: Section 2 will introduce the origins and foundations of responsible AI; Sections 3 will explore each major dimension in detail; Section 4 will identify Future Directions and Recommendations; Section 5 will conclusion of responsible AI.

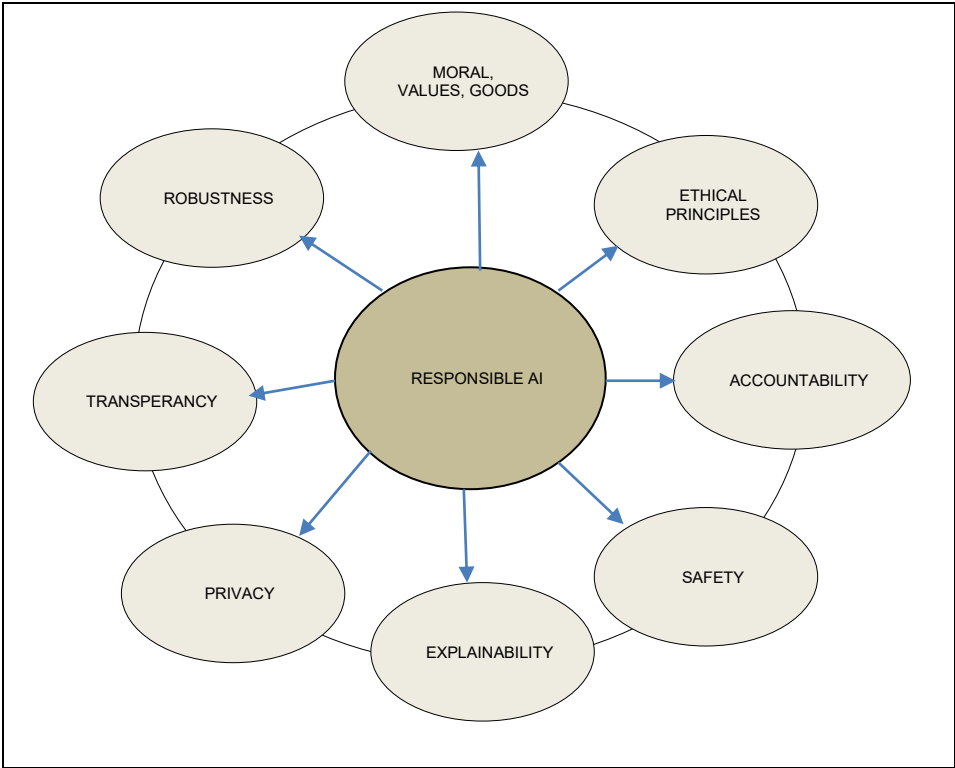


Fig.1. Multidimensional Framework of Responsible AI [5].

2 Background and Theoretical Foundations

Responsible AI has grown considerably over the last decade- a scant attention on algorithmic fairness to a wide sociotechnical paradigm which acknowledges the inter-

action among technical systems, organizational tasks, rules, and social values [8]. The initial research has focused on technological issues like bias detection and its mitigation in machine learning models [9]. Nevertheless, with the introduction of AI systems into the center of high-stakes decision-making, it was obvious that technical solutions did not suffice in regard to countering the ethical, social and governance issues at hand [10]. According to recent studies, the transition to dynamic context-specific and ongoing practice of responsible AI, operationalized through sociotechnical systems instead of being a compliance activity, is advocated instead of passively working through some fixed ethical checklists. This change brings about the necessity of a not only technical but also organizational change, regulatory aid, stakeholder engagement, and continuous oversight and adjustment [11]. The following are the main principles of the modern responsible AI frameworks: fairness, transparency, accountability, privacy, safety, human oversight, and social or environmental well-being. A holistic model proposed by Jiang et al. [1] is the one that covers three interrelated pillars: Intrinsic Security (robustness), Derivative Security (real-world harm mitigation), and Social Ethics (value alignment and accountability) to capture value-additional technical, preventive and ethical aspects of responsible AI [12]. Equally, Gunasekara et al. [5] note that ethical considerations, fairness, transparency and accountability, privacy, safety and human-centered design are important and requires combined methodologies, but not single methods. In the literature, responsible AI is often considered to be multidimensional, which involves coordinated technical, ethical, organizational, and regulatory work [13]. Nevertheless, there are still serious problems in terms of principles operationalization, resolving the conflicting values, and changing the abstract regulations into practice [14]. The subsequent sections will look at these dimensions individually, and generalize existing knowledge and pinpoint important challenges and opportunities. Figure 2 shows evolution of responsible AI (2010–2025). The philosophical grounds of responsible AI systems are consequentialism, deontology, virtue ethics, and care ethics [2], [19]. According to Sridhar [2], fairness, accountability, and transparency are particularly important in situations where AI is used in high-stakes scenarios and this has to be underpinned by existing ethical theories.

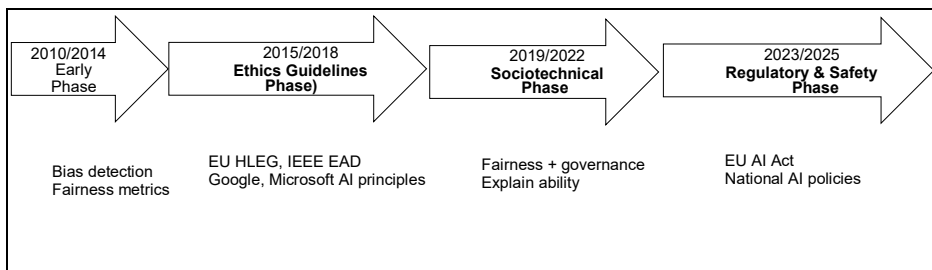


Fig.2. Evolution of Responsible AI (2010–2025) [19]

The key issue, however, is how abstract moral values can be translated in technical and organizational practices [15]. Recent efforts seek to fill this gap with more practical models, including the detailed model by Karras [6], trying to combine ethical the-

ory and technological approaches and focusing on interdisciplinary cooperation between ethicists, technologists, policy makers and domain specialists [2], [15]. Contemporary systems are progressively incorporating equity, openness, responsibility, confidentiality, and security into coherent systems that acknowledge diplomacy and inter-relationships [12], [15], and are dynamic and context-sensitive, as opposed to fixed universal regulations [24]. However, the operationalization of these models is still a challenge, as it is a symptom of consistent discrepancy between ambivalent principles and practical application, as well as, lack of clarity of organizational processes [17], [19]. One of the most reported issues with responsible AI is algorithmic bias, which can be caused by biased training data, the poor model design, or inadequate deployment settings and may occur at any point of the AI lifecycle [18]. There are records of demographic inequalities in facial recognition [19], discrimination in predictive policing [29], and discriminated credit and hiring practices. Technical methods of fairness, like pre-processing, in-processing and post-processing methods, are also in development, such as fairness-conscious algorithms [20] and application-specific fairness requirements in fields such as banking. The biggest current challenge is the need to choose between mathematically incompatible concepts of fairness, including demographic parity and equalized odds, and such decisions are to be made depending on context [28]. Reasonableness in sequential decision-making also finds increasing popularity, whereby new approaches such as the fairness shields of Cano [30] have been implemented to measure cumulative effects with time, but this field is not as advanced as fairness in making a static prediction [21]. Another fundamental problem is transparency and explainability where complex models, especially deep neural networks, are usually not transparent, making it hard to hold people accountable, or to trust them, and fix errors efficiently [22]. Such ambiguity is particularly hazardous in high stakes industries like healthcare, finance and criminal justice areas where the direction of decision-making should be interpretable. True transparency that lets users and regulators challenge and review system decisions is necessary to have a trustworthy AI, as Liu [25] argues.

3 Explainable AI Methods and Tools

To resolve this problem of ambiguity in AI decision-making, scholars have proposed numerous explainable AI (XAI) methods that promote greater levels of transparency and interpretability by both intrinsically interpretable models and post-hoc explanations of complex black-box systems [2]. Explainability is also emphasized by Sridhar [2] as a necessity of responsible AI, since it facilitates human control and responsibility, whereas more recent advances are mentioned by Majjate et al. [15] (attention mechanisms, saliency maps, and counterfactual explanations), Selvam [27] and Sanka [28] demonstrate how XAI tools can be adapted to the needs of a particular domain. Nevertheless, most XAI approaches are incomplete or inaccurate with the most complex models, performance-interpretability trade-offs have not been eliminated, and different forms of explanations must be offered to various stakeholders [29].

In addition to the technical approaches, it has been suggested that transparency is also based on the practices of the organization, and Tabaghdehi et al. [12] suggest implementing a circular model where the transparency and accountability are treated as well as inclusivity, which could involve record-keeping, data and model limitations documentation, and communication with stakeholders [28].

Responsible AI has the principle of accountability and balances between innovation and trust and benefit mechanisms [12] and includes causal responsibility, role responsibility, and legal liability [30], and Cano [30] suggests formal tools to assess agency in autonomous systems. Accountability is still hard to enforce because of the convoluted AI lifecycle, which consists of data providers, developers, deployers, and users, and adaptive system behavior that makes assigning responsibility more challenging [1].

Governance structures thus become very important whereby technical, organizational, and policy elements must be encouraged to foster transparency, accountability, and trust [9] with governance being used as a design principle and not as an auxiliary [16]. The current literature emphasizes multi-level governance (organizational to international) with the help of ethics committees, impact assessments, audits, and stakeholder engagement [23] and states that the use of the static compliance checklists is ineffective in the context of the fast-paced AI environment, which requires dynamic and adaptive forms of governance based on sociotechnical practice [11].

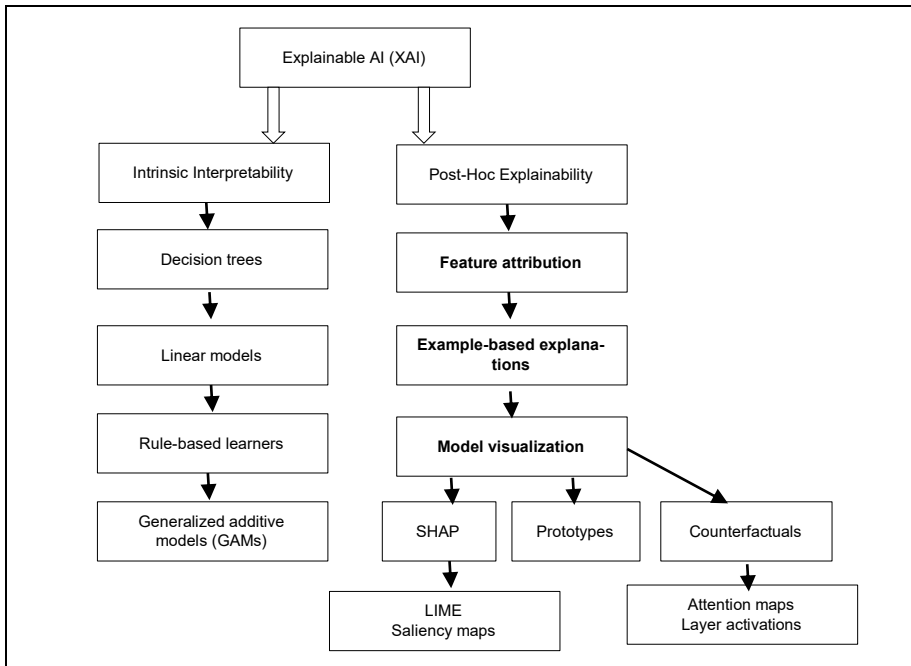


Fig.3. Taxonomy of Explainable AI (XAI) Approaches [26]

Table 1. Limitations of Current XAI Methods [28].

Limitation	Explanation	Example
Misleading saliency maps	The saliency maps are easy to manipulate or it can represent background noise instead of the actual important features.	A heatmap of the irrelevant areas in an image that was classified.
Linear models	The relative change in features can be small and occasionally give a proportionately large change in predictions.	One small variation in the aspects of the input causes a drastic change in the model prediction.
Non-robust explanations	With minor perturbations in the input data, explanations may vary greatly.	The introduction of some random pixels in an image alters the results of the explanation.
SHAP / LIME limitations	In the case of feature attribution, issue might result in volatile or unreliable explanations with regard to the sampling strategy.	SHAP/LIME runs giving slightly dissimilar explanation to the same prediction.
Trade-off: Accuracy vs Interpretability	Deep learning models and models that are highly accurate are not always interpretable, unlike simpler models.	Deep neural network offers higher accuracy, and low explainability in loan approval decisions.

To roll out the responsibility of responsible AI, it is necessary to commit resources, train the workforce, and integrate responsible AI practices into the development processes [8] but interdisciplinary collaboration, which is a critical requirement, is challenged by the use of different disciplinary languages, methods, and priorities among the technologists, ethicists, policymakers, and domain experts [2]. Figure 3 shows taxonomy of explainable AI (XAI) Approaches. Bach et al. [11] present a systematic review of the literature on proposed responsible AI practices in practical environments and find multiple gaps between ideal philosophy and concrete practice. The concept of reliability has expanded in the field of responsible AI to include technical performance, ethical alignment, transparency, accountability, and competence [14]. Trust has been applied not as an individual property but as a result of joint system properties and stakeholder beliefs [17], where the experts are more interested in robustness whereas the end users, in fairness and transparency [24]. Jiang et al. [1] consequently state that credible AI entails the combination of technical integrity and social responsibility, which reflects its conformity to societal norms, which is a noteworthy change in the perception of trust as a technical aspect that used to be applied in the past [23], [24]. The measurement of trustworthiness is still challenging because there are insufficient evaluation processes and restrictions of the accuracy-based measures, which are incapable of considering fairness, robustness, transparency, and value alignment [1], [15]. Reuel et al. [17] suggest a global maturity model that can assist organizations in determining the responsible AI use though creating extensive, fully validated tools is still difficult [11] and researchers point to the necessity of creating evaluation tools that capture actual risks or real-life consequences, stakeholder responses, and system controls instead of laboratory standards [24]. Responsible AI is centred on providing the highest level of safety and robustness given that vulnerabilities, including bias, adversarial attacks, and malfunctioning in distribution shifts may result in misinformation, inequity, breaches of security, physical harm, and destroyed

trust as shown with adversarial examples, emergent behaviour, and fragile performance. In order to enhance robustness, researchers have come up with defensive mechanisms, including safety shields, adversarial training, certified defenses, and ensemble mechanisms. Cano [30] demonstrates how deterministic shielding can benefit real-world applications by keeping systems within safety constraints, especially in high-stakes environments such as autonomous vehicles [1], [14], [30]. Nonetheless, defenses usually do not generalize to new threats, and there is always a gap between the newest attack and the current defenses, which means that emerging attacks are always faster than the newest protection mechanisms and they must be adjusted continuously [1].

4 Future Directions and Recommendations

According to the literature, some of the main research priorities in responsible AI, where the generalization gap in safety and robustness needs to be closed by adaptive defenses, meta-learning and formal verification, able to react to evolving threats, is [1], [14], [30]. The challenge of fairness in sequential decision-making has not been adequately addressed, and new definitions and mitigation measures are needed to consider the cumulative effects [15], [30]. Better explainability is also required since much of the existing XAI methods offer partial or misleading information, requiring faithful, understandable and practical explanations [2], [12], [14], [28]. The other priority is the creation of holistic assessment frameworks, which will quantify fairness, transparency, safety and social impact beyond technical accuracy per se [1], [11], [14], [17]. The literature on policy and governance underlines the necessity of harmonization of international regulations, better accountability and enforcement systems and capacity building of institutions through guidelines, training and communities of practice [1], [9], [11], [16], [18], [26]. Inclusive governance should also be included, which requires the inclusion of the marginalized groups and websites of different values of the society [12], [24], [26]. Last but not least, the interdisciplinary cooperation is viewed as a key to the development of responsible AI in technical, ethical, social and legal spheres with the assistance of common frameworks, institutionalized systems and awareness of various forms of expertise [2], [7], [9], [10], [11], [17], [24], [26].

5 Conclusion

This review summarizes the new 2024/2025 literature on Responsible Artificial Intelligence along nine key dimensions, which are fairness, transparency, accountability, safety, privacy, human-centered design, and regulation. Although significant progress has been achieved in bias reduction, explainability, robustness strategies, privacy-sensitive and other privacy-protecting approaches, as well as novel governance approaches (like the EU AI Act), significant obstacles remain. Some are the generalization gap in AI safety where defenses do not work against current threats [1], weak and limited evaluation protocols that do not consider the real-world effects [14], disjointed

regulations that create inconsistent control [23], and the persistent loophole between principle and application in organizations [11], [17]. The above issues occur because responsible AI is perceived as a checklist, but not as a sociotechnical practice [1], [16]. To help solve them, it is essential to integrate responsible AI into the system lifecycle, build context-dependent strategies, and enhance accountability, interdisciplinary working group, and stakeholder involvement. The way forward in the future should be to develop technical research on fairness in sequential decision-making, robust and generalizable safety techniques, faithful and comprehensible explainability, and holistic assessment frameworks [1], [15], [30], and in the policy realm, regulatory harmonization, greater enforcement, building organizational capacity, and inclusive governance [9], [23], [26]. Finally, accountable AI should guarantee that AI systems benefit human well-being and democratic principles through incorporating technical integrity and social accountability [24], [1], [25], offering a basis of developing AI systems that win and retain people's confidence.

References

- [1] Jiang et al., "Never Compromise to Vulnerabilities: A Comprehensive Survey on AI Governance," arXiv.org, 2025. DOI:[10.48550/arxiv.2508.08789](https://doi.org/10.48550/arxiv.2508.08789)
- Sridhar, "Ethical frameworks for responsible ai development: Challenges and implementation strategies," *World Journal of Advanced Engineering Technology and Sciences*, 2025. DOI: [10.30574/wjaets.2025.15.1.0420](https://doi.org/10.30574/wjaets.2025.15.1.0420)
- Firmansyah et al., "The Future of Ethical AI," *Advances in computational intelligence and robotics book series*, 2024. DOI: [10.4018/979-8-3693-3860-5.ch005](https://doi.org/10.4018/979-8-3693-3860-5.ch005)
- Bach et al., "Insights into suggested Responsible AI (RAI) practices in real-world settings: A systematic literature review," 2024. DOI: [10.21203/rs.3.rs-4707426/v1](https://doi.org/10.21203/rs.3.rs-4707426/v1)
- Gunasekara et al., "A Systematic Review of Responsible Artificial Intelligence Principles and Practice," *Applied system innovation*, 2025. DOI: [10.3390/asi8040097](https://doi.org/10.3390/asi8040097)
- Karras, "On Modelling a Reliable Framework for Responsible and Ethical AI in Digitalization and Automation: Advancements and Challenges," *Financial Engineering*, 2025. DOI: [10.37394/232032.2025.3.29](https://doi.org/10.37394/232032.2025.3.29)
- Tiwari, "Navigating the Frontier: Theoretical Frameworks and Technical Approaches to Responsible AI," *Indonesian Journal of Computer Science*, 2025. DOI: [10.33022/ijcs.v14i2.4789](https://doi.org/10.33022/ijcs.v14i2.4789)
- Shehu et al., "Ethical and Responsible AI in Engineering and Construction Projects: Governance, Trust, and Human-Centered Design," *Scientific journal of engineering and technology*, 2025. DOI: [10.69739/sjet.v2i2.833](https://doi.org/10.69739/sjet.v2i2.833)
- Ismail et al., "Ethical and Governance Frameworks for Artificial Intelligence: A Systematic Literature Review," *International journal of interactive mobile technologies*, 2025. DOI: [10.3991/ijim.v19i14.56981](https://doi.org/10.3991/ijim.v19i14.56981)
- Memarian et al., "A review of caveats and future considerations of research on responsible AI across disciplines," *Human Technology*, 2024. DOI: [10.14254/1795-6889.2024.20-3.4](https://doi.org/10.14254/1795-6889.2024.20-3.4)

10. Bach et al., "Insights into suggested Responsible AI (RAI) practices in real-world settings: a systematic literature review," *AI and ethics*, 2025. DOI: [10.1007/s43681-024-00648-7](https://doi.org/10.1007/s43681-024-00648-7)
11. Tabaghdehi et al., "AI ethics in action: a circular model for transparency, accountability and inclusivity," *Journal of Managerial Psychology*, 2025. DOI: [10.1108/jmp-03-2024-0177](https://doi.org/10.1108/jmp-03-2024-0177)
12. "Principles and Methods for Transparency, Interpretability, Trust, Accountability, and Responsibility in AI-Based Predictive Modelling," *Computer Science, Engineering and Technology*, 2025. DOI: [10.46632/cset/3/3/10](https://doi.org/10.46632/cset/3/3/10)
13. Kowald et al., "Establishing and Evaluating Trustworthy AI: Overview and Research Challenges," 2024. DOI: [10.48550/arxiv.2411.09973](https://doi.org/10.48550/arxiv.2411.09973)
14. Majjate et al., "Advancing Ethical Ai: Emerging Trends in Transparency, Fairness, and Explainability," 2025. DOI: [10.1109/iccsc66714.2025.11134950](https://doi.org/10.1109/iccsc66714.2025.11134950)
15. Sistla, "AI with Integrity: The Necessity of Responsible AI Governance," 2024. DOI: [10.47363/jaicc/2024(3)e180](https://doi.org/10.47363/jaicc/2024(3)e180)
16. Reuel et al., "Responsible AI in the Global Context: Maturity Model and Survey," 2024. DOI: [10.48550/arxiv.2410.09985](https://doi.org/10.48550/arxiv.2410.09985)
17. Tiwari et al., "Responsible AI in Action," 2025. DOI: [10.4018/979-8-3373-0877-7.ch008](https://doi.org/10.4018/979-8-3373-0877-7.ch008)
18. Siddiqui, "A Comprehensive Review of AI: Ethical Frameworks, Challenges, and Development," *Adhyayan: A Journal of Management Sciences*, 2024. DOI: [10.21567/adhyayan.v14i1.12](https://doi.org/10.21567/adhyayan.v14i1.12)
19. Jadaan et al., "Ethics in AI and Computation in Automated Decision-Making," *Advances in computational intelligence and robotics book series*, 2024. DOI: [10.4018/979-8-3693-6230-3.ch012](https://doi.org/10.4018/979-8-3693-6230-3.ch012)
20. Naripeddy et al., "Ethical Challenges and Solutions in the Deployment of Artificial Intelligence Sys-tems," *Deleted Journal*, 2025. DOI: [10.60087/jaigs.v8i02.392](https://doi.org/10.60087/jaigs.v8i02.392)
21. Vemulapalli, "Charting the Ethical Horizon," 2024. DOI: [10.1201/9781003497189-16](https://doi.org/10.1201/9781003497189-16)
22. Kumar et al., "Strengthening AI governance: International policy frameworks, security challenges, and ethical AI deployment," *ITU journal*, 2025. DOI: [10.52953/cogc3309](https://doi.org/10.52953/cogc3309)
23. Ibitoye et al., "Rethinking responsible AI from ethical pillars to sociotechnical practice," *AI and ethics*, 2025. DOI: [10.1007/s43681-025-00809-2](https://doi.org/10.1007/s43681-025-00809-2)
24. Liu, "Building Trustworthy AI: Transparency, Fairness, and Governance in the Digital Age," 2025. DOI: [10.63802/afs.v1.i1.80](https://doi.org/10.63802/afs.v1.i1.80)
25. Nangoy et al., "Toward Ethical AI: Strategies for Responsible AI Governance," *Journal of business and management studies*, 2025. DOI: [10.32996/jbms.2025.7.5.13](https://doi.org/10.32996/jbms.2025.7.5.13)
26. Selvam, "Ethical AI for Personalized Banking: Addressing Bias and Fairness Challenges," *LatIA*, 2025. DOI: [10.62486/latia2025361](https://doi.org/10.62486/latia2025361)
27. Sanka, "Ethical AI in Big Data: Challenges in Bias, Fairness, and Transparency," *International Journal of Scientific Research in Science and Technology*, 2025. DOI: [10.32628/ijrst25121220](https://doi.org/10.32628/ijrst25121220)

28. Kasoju et al., "The ethics of AI decision-making: Balancing innovation and accountability," *International Journal of Science and Research Archive*, 2024. DOI: [10.30574/ijrsra.2024.12.2.1548](https://doi.org/10.30574/ijrsra.2024.12.2.1548)
29. Cano, "Towards Responsible AI: Advances in Safety, Fairness, and Accountability of Autonomous Systems," *arXiv.org*, 2025. DOI: [10.48550/arxiv.2506.10192](https://doi.org/10.48550/arxiv.2506.10192)

Open Access This chapter is licensed under the terms of the Creative Commons Attribution-NonCommercial 4.0 International License (<http://creativecommons.org/licenses/by-nc/4.0/>), which permits any noncommercial use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

