



International Journal of Research Publication and Reviews

Journal homepage: www.ijrpr.com ISSN 2582-7421

Artificial Intelligence (AI) Misuse and the Law: Addressing Societal Risks in the Digital Age

Chester L. Cofino

College of Computer Studies, Silliman University
chesterlcofino@su.edu.ph

ABSTRACT

Artificial Intelligence (AI) has become a game-changer in this digital age, changing how online transactions work, from governments and private services to how people respond to their daily routines. However, when this new technology is improperly used, it will present significant ethical and legal concerns. This study checks on the misuse of AI tools, the foundation of AI based on its conceptual and theoretical aspects, and the risks they pose to society. The study further examines how different countries' legal and regulatory frameworks protect their community from these difficulties, emphasizing the advantages and disadvantages of existing governance structures of the data protection laws, cybersecurity rules, and ethical AI principles that some countries offer protection through. To solve these problems, humans need to work together and protect each other from the misuse of AI. The paper argues that societies are susceptible to the disruptive effects of AI misuse without a strong and comprehensive regulatory policy adaptable to international and local environments.

Keywords: Artificial Intelligence Misuse, Societal Risks, Algorithmic Bias, Deepfakes and Misinformation, Legal and Regulatory Frameworks

Introduction

Artificial Intelligence (AI) has quickly become one of today's most important technologies (Stamova & Draganov, 2020). It is changing the industries we worked in years before, making daily transactions of each individual more efficient, and improving decision-making (Prasanth et al., 2023; Stone et al., 2020). Its incorporation into society promises advancement, efficiency, and innovation (Ali, 2024). But with these "big bang" improvements comes a growing worry about how these AIs may cause bad things for society. Criminals and actors today are also using AI more and more to gain profit through committing fraud, creating fake identities, automating cyberattacks, and changing how information systems work (King et al., 2020; Treleaven et al., 2023). The spread of AI-generated false information is one of the most serious threats because it erodes trust in the media, disrupts democratic processes, and sways public opinion with false stories (Sophia LI, 2025).

We can't ignore the risks these actors and other lawbreakers pose to society and how they use AI improperly. According to (Sholademi, 2024), common attacks and actions these people do are deepfakes, AI-generated phishing attempts, and automated chatbots made for scams that show how technology can lie and take advantage of people when it is misused. These problems are getting serious because there aren't enough clear legal frameworks to protect society against these attacks. This is because current policies are often outdated and can't keep up with how quickly technology changes (Chubb et al., 2022). If nothing is done, societies could lose public trust, become more open to cybercrime, and lose the integrity of their digital systems.

To solve these problems against AI misuse, the world needs more than tech fixes and bright minds. We also need strong legal and ethical frameworks that put accountability, openness, and responsible use of AI. Countries worldwide, through their governments and private organizations, can reduce risks and encourage innovation by creating flexible policies and strengthening regulatory systems. Thus, this research needs policy-driven approaches to combat AI misuse and misinformation, thereby protecting individual rights and ensuring societal stability.

Conceptual and Theoretical Foundations

AI misuse derives from its dual-use characteristic, wherein systems intended for innovation and advancement may be readily adapted for evil or exploitative purposes (Pöhler et al., 2024). According to (Mittra et al., 2025) the growth of the different types of artificial intelligence, such as machine learning, natural language processing, and generative models, has surpassed the development of legal, ethical, and institutional protections, which has led to chances for misuse. This problem is escalated because not everyone has access to AI experts, the public is not well-informed, and there is not enough regulatory control (Smuha, 2021). This resulted in bad people using AI for fraud, spreading false information, abusing surveillance, and other negative activities (Bazarkina & Pashentsev, 2020; Blauth et al., 2022). Without proactive measures, the misuse of AI increases risks to society by harming the public trust, weakening democratic institutions, and worsening digital-age disparities.

Understanding how AI misuse and the risks of society require grounding in theoretical perspectives that explain the relationship between technology, law, and culture(Christou, 2025). Ethics in the use of technology provides a good foundation to examine how values such as fairness, accountability, and human dignity should help develop and deploy AI(Ryan & Stahl, 2021). The perspective of (Akhundov, 2025)emphasizes that technology is never neutral; rather, it embodies the intentions and biases of its designers, making ethical oversight essential to prevent misuse.

From a legal perspective, the discussion of some scholars and lawmakers between legal positivism and natural law offers conflicting views through which AI misuse can be addressed. Grumulaitis(Grumulaitis, 2025)states that legal positivism views law as a product of formal rules and institutions, implying that regulating AI requires codified statutes, policies, and enforcement mechanisms regardless of moral considerations. In contrast, natural law theory stresses that laws should reflect general ethical principles and human rights, suggesting that AI regulation must exceed the requirements, like identifying the best practices to protect human dignity, justice, and the well-being of society(Chaturvedi, 2025). This tension underscores the difficulty of creating AI policies that are both enforceable and ethically sound.

The debate between legal positivism and natural law further influences governance approaches. The book of (Jha & Pandey, 2023)highlighted how legal positivism supports the importance of legal policy, regulatory bodies, and guidelines in compliance with these legal requirements to hold bad people accountable. At the same time, natural law perspectives stress that AI regulation must also reflect the fairness of justice and human rights (Hakan Kan, 2024).

AI misuse demonstrates that technology cannot be separated from the ethical, legal, and societal contexts in which it operates(Anderljung et al., 2025; Pöhler et al., 2024). Perspectives from technology ethics, legal theory, and risk society highlight the complex nature of AI misuse, showing that it is both a moral and regulatory challenge and a systemic risk in modern societies.

AI Misuse in Practice: Societal Risks

The foundations of AI provide a clear insight into how people misuse the tools, harming society in ethical and legal dimensions. As a result, injuries among the victims of misuse of these technologies raised concern in everyday life.AI technologies are increasingly exploited in ways that harm privacy, distort information, promote discrimination, and threaten security, hence compounding societal hazards in the digital era.Examining these real-world situations is vital to illustrate the tangible repercussions of misuse and to underscore the urgent need for effective legal and governance systems.

Data Privacy and Surveillance

AI-driven technologies have escalated concerns about data privacy and surveillance(Singh, 2024). Why did this happen? This is because of the massive data collection using the internet, facial recognition, and predictive analytics(Maphosa, 2024). Most governments and corporations worldwide accept using AI through their systems to monitor individuals, track habits (Cihon et al., 2021), and collect personal information, frequently without explicit agreement (Longpre et al., 2024). Systems with facial recognition technology are now used in public settings, law enforcement, and border control. With these tools used in recognizing personal information, worries about ongoing monitoring and the erosion of anonymity (Neroni Rezende, 2022; Solarova et al., 2023). The study of (Limantè, 2024)further shows that these technologies are sensitive to prejudice; if there is a case of misidentification, these will disproportionately harm the individual, heightening the dangers of discrimination and unjust profiling. Such behaviors illustrate the growing contradiction between technological efficiency and preserving civil liberties.

There are some effects of AI spying that extend beyond individual privacy. According to(Büchi et al., 2022), the continued use of the AI spying tools in monitoring generated discouragement in some individuals, who limited how they interacted with others in society, being afraid that they were being monitored. At the same time, companies use user data for targeted advertising, consumer profiling, and behavioral prediction, tactics that erode autonomy and subject individuals to manipulation.

One key difficulty in real-world scenarios is that governments and law enforcement agencies use facial recognition technology (FRT). In certain countries, AI-powered surveillance cameras scan public locations to track movements and identify(Fontes et al., 2022). While proponents argue this enhances security, numerous reports reveal wrongful arrests and racial misidentification, particularly in the U.S. and U.K.(Al-Dulaimi& Mohammed, 2025; McMullen, 2020), raising issues about bias, discrimination, and violations of the assumption of innocence. The lack of openness in how data is collected and maintained further exacerbates privacy concerns, as citizens typically remain unaware of the level of monitoring (Fabrègue&Bogoni, 2023).

In the corporate sector, widespread data exploitation by digital companies offers another urgent concern (Lindman et al., 2023). This is because the platforms they provide to society frequently collect and analyze user data with AI-driven algorithms for profit purposes(Rainy, 2025). These approaches fuel targeted advertising and predictive analytics, yet often occur without significant consumer agreement. An example of this is the issue about the Cambridge Analytica incident that highlighted how data acquired from millions of users might be exploited to affect political outcomes(Boerboom&Boerboom Lee, 2020), illustrating the societal risks of unrestrained corporate spying.

Another critical issue is the absence of robust legal safeguards in many jurisdictions, allowing governments and corporations to operate in legal grey areas(Babikian, 2023). Countries with weak or outdated data protection laws fail to protect citizens from excessive surveillance, leaving personal data vulnerable to misuse, breaches, or authoritarian overreach (Bentotahewa et al., 2022).

Algorithmic Bias and Discrimination

Even though these AI systems are designed to be neutral(Shalevska& Walker, 2025), they sometimes create and even cause social inequalities(Zajko, 2022). For example, in hiring, there are recruitment applications that the algorithms are trained to disadvantage a certain attribute of the applicant, reinforcing long-standing labor market disparities(Njoto et al., 2022). Another is in law enforcement, where predictive policing tools have been criticized for targeting communities based on color, resulting in bias in terms of physical features (Bates, 2024; Yen & Hung, 2021), as these systems rely on historical crime data that reflect biased policing practices. In healthcare, several AI diagnostic systems have shown lower accuracy for patients from marginalized groups, resulting in unfair treatment recommendations (Owolabi et al., 2025; Seyyed-Kalantari et al., 2021). These examples show that AI applications aren't just theoretical problems but real things that affect society.

These biases arise because of how datasets are trained and installed in AI systems for decision-making processes that reflect an imbalanced society(Lainjo, 2023). Because of this, discrimination based on algorithms implemented in AI systems makes people less trusting in AI, but it also raises important moral and legal problems about justice, responsibility, and the preservation of human rights.

Some studies highlight fairness, transparency, and accountability as foundational principles for mitigating algorithmic bias and discrimination. Fairness involves designing systems that prevent inequities and ensure equal treatment across different demographic groups(Pulivarthy& Whig, 2024). Transparency requires that the decision-making processes of AI systems be explainable and transparent to stakeholders to prevent scrutiny and oversight. Accountability, on the other hand, means assigning someone that is responsible when AI systems hurt people, whether that person is a creator, an institution that uses the technology, or a regulator (Cheong, 2024). Together, these ideas make up the moral foundation of responsible AI governance.

Deepfakes and Misinformation

Tools that were developed by companies or individuals with AI integration for media-related services, especially deepfakes, are used to spread false information to some communication tools like social media and news platforms(Veeramy& Pieterse, 2022), which has caused a serious problem in politics(Battista, 2024), security, and public trust. Some images, audio, or videos seen on the internet are not trusted anymore(George &Hovan George, 2023).It is really disturbing that there are multimedia materials that look very real and can make it look like people are saying or doing things they never did. Studies have shown that these technologies have been used to hurt the reputations of politicians, spread false stories, and change people's minds. For example, videos generated using AI were spread by some individuals during election campaigns to share misinformation or to hurt candidates' credibility, lower voter confidence, and divide societies(Islam et al., 2024). In addition, some AI-generated media has been used in security situations, such as scams and identity theft(Sholademi, 2024), making things more dangerous for people and organizations.

The dissemination of this deepfake content spread very quickly as other users usually share the content with each of their connected users. This caused a hard time detecting and mitigating using verification tools for fact-checking. Fact-checking methods that have been around for a long time often have difficulty keeping up with the speed and complexity of AI-generated false information(Taiwo, 2025). Also, once this information is shared on social media, it spreads quickly and stays available even after it is wrong. This has a long-term effect on how people see things. People can't tell the difference between real news and fake news as easily when they don't trust the government. This makes political processes less stable.

Deepfake content and other AI-generated false information caused damage or had a big impact on individuals and society in terms of democracy. Because democratic processes rely on how citizens participate and are informed, the proliferation of fabricated AI-generated content(Kreps &Kriner, 2023) may alter electoral outcomes, exacerbate societal divisions, and undermine institutional legitimacy.

Another concerning issue about misinformation is called false stories that usually target minorities, making divisive issues that can make social tensions worse, encourage violence, and weaken social ties(Khosa&Abdulkareem, 2023). The term "liar's dividend," which means that real content is thrown out as fake because of deepfakes (Gondaliya, 2025), makes people even less likely to trust information. This problem proved how important it is to have robust rules and a set of programs that teach people how to use technology to find and stop false information made by AI to protect democracy and social stability.

Cybersecurity and Autonomous Systems

Aside from misinformation and algorithmic discrimination, actors go beyond what these tools distract society from, like more aggressive cyberattacks and the potential deployment of the so-called autonomous weapons. These assaults take advantage of weaknesses with an accuracy that has never been seen before(Akhtar &TajbiulRawol, 2024). They automate phishing attempts to target people (Guembe et al., 2022), break passwords(Gürfidan et al., 2023), and get into systems (Prince et al., 2024). The newest technology with AI-enhanced capabilities(Ali Syed, 2025)may make phony emails, voice recordings, or deepfake movies that trick people and surpass standard security procedures. Autonomous weapons are a growing security danger outside of cyberspace (Altmann, 2019). AI systems used in military technologies could make life-or-death decisions without direct human oversight(Špelda, 2024). Scholars and policymakers caution that these advancements obscure the distinction between technical innovation and global security issues.

The research presented by (Haskard & Herath, 2025) indicated that if these cyberattacks and autonomous weapons are integrated to attack a specific organization or individual, this creates a general issue of trust and control in this digital age. The dangers of autonomous weapons worsen these concerns since they take human judgment out of conflict situations, which is both a tactical and an ethical problem.

The development and free use of high-risk AI applications have led to more discussions about who would be the responsible individual or organization to address the problem and what standard rules to use in fighting this problem. One important question is who should be accountable when AI systems cause harm: the people who made the algorithms, the companies that employ them, or the governments that don't control their use. The information on these AI-powered cyberattacks done by bad actors is hidden worldwide. Autonomous weapons also raise problems about who is responsible for what happens in battle, especially if robots make mistakes that hurt civilians or break international humanitarian law (Selbst et al., 2020).

Local and international discussions between the public and private sectors were conducted to emphasize the need for a proactive framework to mitigate these concerns. Many proposals include international treaties to ban or limit autonomous weaponry, stricter cybersecurity measures, and liability laws suited to AI misuse. Also, debates draw attention to the need to foster innovation and maintain safety, since overly stringent rules could inhibit technological progress, while permissive policies may expose societies to catastrophic misuse.

Lastly, AI misuse really violates data privacy, discrimination, deepfake misinformation, cyberattacks, and autonomous weaponry, displaying the hazards it poses to society. These activities endanger individual rights and security and threaten democratic institutions, social stability, and world peace. As the actual realities of AI misuse continue to emerge, they underscore the critical need for robust legal, ethical, and governance frameworks that can secure society while encouraging innovation.

Legal and Regulatory Perspectives

Society's concern about the misuse of AI needs to be considered by government officials in every country, and they need to create clear regulations and establish them properly. Legal and regulatory perspectives provide the framework for societies to balance innovation with accountability, ensuring that AI technologies are used responsibly while respecting individual rights and public interests. To do this, existing laws, ethical principles, and policy debates are crucial, and it investigates how legal frameworks respond to the issues of AI misuse and identifies the gaps that must be addressed to develop effective and adaptive governance systems in the digital era.

Existing Legal Frameworks

Some existing legal frameworks for AI misuse largely focus on neighboring topics such as data protection, cybersecurity, consumer protection, and human rights, rather than AI-specific law (Ijaiya & Odumuwagun, 2024). For example, the European Union's General Data Protection Regulation (GDPR) sets global standards for personal data handling, demanding transparency and user consent in data processing concepts that are extremely applicable to AI systems (de Magalhães, 2020). Other jurisdictions, including Singapore, Canada, and the Philippines, have established national data privacy legislation that offers basic rights against monitoring and unlawful data usage (Corning, 2024; Kennedy et al., 2009). While these regulations provide some safeguards, these frameworks often cannot be updated easily to keep pace with the rapid advancement of AI technology and its misuse in areas like deepfakes, algorithmic bias, and autonomous systems.

Ethical AI guidelines have been produced by organizations such as the UNESCO, and the European Commission, emphasizing values of justice, transparency, and accountability. However, these rules are mostly voluntary and non-binding, limiting their efficacy in reducing harmful applications.

Giant countries around the world have already initiated the crafting of their framework to combat AI misuse. First, the European Union has already adopted and implemented its AI Act, proposing to build the first complete regulatory framework that categorizes AI systems by risk category and sets rigorous constraints on high-risk applications (Cabrera et al., 2025). In contrast, the United States uses a more market-driven strategy, emphasizing innovation and self-regulation, which critics contend leaves huge accountability gaps. Meanwhile, China has integrated AI governance into its broader state control structure, focusing on censorship, national security (Zeng, 2020), and maintaining societal stability. These diverse approaches underline the problem of obtaining a worldwide consensus on AI governance.

Challenges in Regulation

In AI governance, one challenge that concerns any policy implementor is the jurisdictional gaps that occur from the global nature of AI development and deployment. AI-driven systems, particularly in data processing, deepfake generation, and cyberattacks, often operate across numerous jurisdictions simultaneously. Local laws in every country are naturally limited in scope, focusing on the local level aligned to the needs of its constituents, leaving governments unable to oversee acts that originate outside yet have domestic implications. For example, a bad actor operating in one nation can conduct AI-powered misinformation campaigns or cyber breaches, destabilizing institutions in another. However, legal remedies are typically impossible due to sovereignty restrictions. This lack of legal clarity creates gaps that allow malevolent AI use to flourish in a cross-border environment.

Even when legal frameworks exist, enforcement and cross-border accountability remain poor. Usually, governance frameworks for AI rely on voluntary ethical norms rather than formal international agreements, making compliance inconsistent. Enforcement authorities also suffer from AI misuse, such as automated cyberattacks or autonomous decision-making systems, when identifying the accountable individual is complicated. Differing regulatory

philosophies further impede international cooperation. While the EU stresses precautionary regulation through initiatives like the AI Act (Hacker, n.d.), other jurisdictions promote innovation or state control, leading to fragmented standards. Without unified norms and cross-border enforcement tools, societies confront chronic weaknesses in tackling AI misuse that crosses national boundaries.

Ethical and Human Rights Considerations

AI governance emphasizes that artificial intelligence development, deployment, and regulation must prioritize protecting fundamental rights such as privacy, freedom of expression, equality, and non-discrimination. Unlike purely technical or market-driven perspectives, this approach positions human dignity and autonomy at the core of AI governance. It recognizes that technologies like facial recognition, algorithmic decision-making, and AI-driven surveillance can disproportionately harm vulnerable groups, erode civil liberties, and amplify systemic inequalities if left unchecked.

Conclusion

Artificial Intelligence offers large opportunities for innovation, but its misuse in illegal activities poses serious risks to trust, security, and social stability. Combating these challenges requires cooperation with society. Creating a standardized and comprehensive legal framework to fight against its misuses, protect individuals and the community. Technology companies and digital platforms worldwide must implement responsible AI practices, deploy detection tools, and promote transparency to prevent malicious use. At the same time, society must be alert and improve digital literacy and critical thinking skills to fight against AI-driven lies. When these initiatives unite, society may find a balance between supporting innovation and protecting against misuse. This will make sure that AI stays a tool for advancement instead of being used to exploit people in the digital era.

Recommendation

To reduce the threats that AI misuse and misinformation pose to society, it is suggested that governments create flexible and inclusive AI rules that keep up with technological changes. They should also set up special task forces to investigate and monitor AI-related crimes. The private sector should create its rules, make AI development more open, and utilize powerful detection techniques to stop people from using their platforms for bad things. In addition, schools, news organizations, and community groups need to work together to promote digital literacy programs that teach people how to spot and fight AI-driven lies. Finally, encouraging international cooperation and exchanging information will help ensure that worldwide standards are the same. This will make it tougher for bad people to exploit legal or technological loopholes. AI can be guided toward responsible and ethical use by combining efforts from the legal, governmental, corporate, and social sectors. This will help keep trust and security in the digital age.

References

- Akhtar, Z. Bin, & TajbiulRawol, A. (2024). Enhancing Cybersecurity through AI-Powered Security Mechanisms. *IT Journal Research and Development*, 9(1), 50–67. <https://doi.org/10.25299/itjrd.2024.16852>
- Akhundov, A. (2025). The Role of Ethics in Modern Technology Development. *Porta Universorum*, 1(4), 169–177. <https://doi.org/10.69760/portuni.0104017>
- Al-Dulaimi, A. O. M., & Mohammed, M. A.-A. W. (2025). Legal responsibility for errors caused by artificial intelligence (AI) in the public sector. *International Journal of Law and Management*. <https://doi.org/10.1108/IJLMA-08-2024-0295>
- Ali, R. (2024). Spectrum of Research and Reviews Article Info. *Spectrum of Research and Reviews*, 1(2), 79–91. <https://thesrr.net>
- Ali Syed, S. (2025). Adversarial AI and Cybersecurity: Defending Against AI-Powered Cyber Threats.
- Altmann, J. (2019). Autonomous Weapon Systems – Dangers and Need for an International Prohibition (pp. 1–17). https://doi.org/10.1007/978-3-030-30179-8_1
- Anderljung, M., Hazell, J., & von Knebel, M. (2025). Protecting society from AI misuse: when are restrictions on capabilities warranted? *AI & SOCIETY*, 40(5), 3841–3857. <https://doi.org/10.1007/s00146-024-02130-8>
- Babikian, J. (2023). Law Research Journal Securing Rights: Legal Frameworks for Privacy and Data Protection in the Digital Era. *Law Research Journal*, 1(2).
- Bates, T. (2024). Technology and Culture: How Predictive Policing Harmfully Profiles Marginalized People Groups. *California Sociology Forum*, 6(1), 18–27.
- Battista, D. (2024). Political communication in the age of artificial intelligence: an overview of deepfakes and their implications. *Society Register*, 8(2), 7–24. <https://doi.org/10.14746/sr.2024.8.2.01>
- Bazarkina, D. Yu., & Pashentsev, E. N. (2020). Malicious Use of Artificial Intelligence. *Russia in Global Affairs*, 18(4), 154–177. <https://doi.org/10.31278/1810-6374-2020-18-4-154-177>

- Bentotahewa, V., Hewage, C., & Williams, J. (2022). The Normative Power of the GDPR: A Case Study of Data Protection Laws of South Asian Countries. *SN Computer Science*, 3(3), 183. <https://doi.org/10.1007/s42979-022-01079-z>
- Blauth, T. F., Gstrein, O. J., & Zwitter, A. (2022). Artificial Intelligence Crime: An Overview of Malicious Use and Abuse of AI. *IEEE Access*, 10, 77110–77122. <https://doi.org/10.1109/ACCESS.2022.3191790>
- Boerboom, C., & Boerboom Lee, C. (2020). Cambridge Analytica: The Scandal on Data Privacy. <https://digitalcommons.augustana.edu/ethicscontest/18>
- Büchi, M., Festic, N., & Latzer, M. (2022). The Chilling Effects of Digital Dataveillance: A Theoretical Model and an Empirical Research Agenda. *Big Data & Society*, 9(1). <https://doi.org/10.1177/20539517211065368>
- Cabrera, B. M., Luiz, L. E., & Teixeira, J. P. (2025). The Artificial Intelligence Act: Insights regarding its application and implications. *Procedia Computer Science*, 256, 230–237. <https://doi.org/10.1016/j.procs.2025.02.116>
- Chaturvedi, M. A. (2025). AI And Ethics: Jurisprudential Psychology To Investigate How Far It Can Be Regulated. *American Journal of Psychiatric Rehabilitation*, 784–796. <https://doi.org/10.69980/ajpr.v28i1.194>
- Cheong, B. C. (2024). Transparency and accountability in AI systems: safeguarding wellbeing in the age of algorithmic decision-making. *Frontiers in Human Dynamics*, 6. <https://doi.org/10.3389/fhumd.2024.1421273>
- Christou, P. A. (2025). A critical inquiry into the personal and societal perils of Artificial Intelligence. *AI and Ethics*, 5(3), 2547–2555. <https://doi.org/10.1007/s43681-024-00556-w>
- Chubb, J., Cowling, P., & Reed, D. (2022). Speeding up to keep up: exploring the use of AI in the research process. *AI & SOCIETY*, 37(4), 1439–1457. <https://doi.org/10.1007/s00146-021-01259-0>
- Cihon, P., Schuett, J., & Baum, S. D. (2021). Corporate Governance of Artificial Intelligence in the Public Interest. *Information*, 12(7), 275. <https://doi.org/10.3390/info12070275>
- Coming, G. P. (2024). The diffusion of data privacy laws in Southeast Asia: learning and the extraterritorial reach of the EU's GDPR. *Contemporary Politics*, 30(5), 656–677. <https://doi.org/10.1080/13569775.2024.2310220>
- de Magalhães, S. T. (2020). The European Union's General Data Protection Regulation (GDPR). In *Cyber Security Practitioner's Guide* (pp. 529–558). WORLD SCIENTIFIC. https://doi.org/10.1142/9789811204463_0015
- Fabrègue, B. F. G., & Bogoni, A. (2023). Privacy and Security Concerns in the Smart City. *Smart Cities*, 6(1), 586–613. <https://doi.org/10.3390/smartcities6010027>
- Fontes, C., Hohma, E., Corrigan, C. C., & Lütge, C. (2022). AI-powered public surveillance systems: why we (might) need them and how we want them. *Technology in Society*, 71, 102137. <https://doi.org/10.1016/j.techsoc.2022.102137>
- George, A. S., & Hovan George, A. S. (2023). Deepfakes: The Evolution of Hyper realistic Media Manipulation. *Partners Universal Innovative Research Publication*, 1(2). <https://doi.org/10.5281/zenodo.10148558>
- Gondaliya, H. (2025). The Rise of AI-Generated Deepfakes: Techniques, Challenges, Implications, and Countermeasures.
- Grumulaitis, A. (2025). Legal Regulation of AI and Morality: The Artificial Intelligence Act in the Context of Natural Law and Legal Positivism. *Teisė*, 134, 27–47. <https://doi.org/10.15388/Teise.2025.134.3>
- Guembe, B., Azeta, A., Misra, S., Osamor, V. C., Fernandez-Sanz, L., & Pospelova, V. (2022). The Emerging Threat of Ai-driven Cyber Attacks: A Review. *Applied Artificial Intelligence*, 36(1). <https://doi.org/10.1080/08839514.2022.2037254>
- Gürfidan, R., Ersoy, M., & Kilim, O. (2023). AI-Powered Cyber Attacks Threats and Measures (pp. 434–444). https://doi.org/10.1007/978-3-031-31956-3_37
- Hacker, P. (n.d.). AI Regulation in Europe: From the AI Act to Future Regulatory Challenges. <https://carnegieendowment.org/2023/02/14/lessons-from-world-s-two-experiments-in-ai->
- Hakan Kan, C. (2024). ARTIFICIAL INTELLIGENCE (AI) IN THE AGE OF DEMOCRACY AND HUMAN RIGHTS: NORMATIVE. *International Journal of Eurasian Education and Culture*. <https://doi.org/10.35826/ijoecc.1825>
- Haskard, A., & Herath, D. (2025). Secure Robotics: Navigating Challenges at the Nexus of Safety, Trust, and Cybersecurity in Cyber-Physical Systems. *ACM Computing Surveys*, 57(9), 1–48. <https://doi.org/10.1145/3723050>
- Ijaiya, H., & Odumuwagon, O. O. (2024). Advancing Artificial Intelligence and Safeguarding Data Privacy: A Comparative Study of EU and US Regulatory Frameworks Amid Emerging Cyber Threats. *International Journal of Research Publication and Reviews*, 5(12), 3357–3375. <https://doi.org/10.55248/gengpi.5.1224.250110>

- Islam, M. B. E., Haseeb, M., Batool, H., Ahtasham, N., & Muhammad, Z. (2024). AI Threats to Politics, Elections, and Democracy: A Blockchain-Based Deepfake Authenticity Verification Framework. *Blockchains*, 2(4), 458–481. <https://doi.org/10.3390/blockchains2040020>
- Jha, A., & Pandey, M. (2023). PRINCIPLES AND THEORIES OF HUMAN RIGHTS. In PRINCIPLES AND THEORIES OF HUMAN RIGHTS. Dominant Publishers & Distributors Pvt Ltd.
- Kennedy, G., Doyle, S., & Lui, B. (2009). Data protection in the Asia-Pacific region. *Computer Law & Security Review*, 25(1), 59–68. <https://doi.org/10.1016/j.clsr.2008.11.006>
- Khosa, D., & Abdulkareem, K. (2023). Breaking the cycle: Presenting insights and strategies to overcome violent conflicts hindering social cohesion and progress in South African communities. *Global Change, Peace & Security*, 35(2), 161–184. <https://doi.org/10.1080/14781158.2024.2407827>
- King, T. C., Aggarwal, N., Taddeo, M., & Floridi, L. (2020). Artificial Intelligence Crime: An Interdisciplinary Analysis of Foreseeable Threats and Solutions. *Science and Engineering Ethics*, 26(1), 89–120. <https://doi.org/10.1007/s11948-018-00081-0>
- Kreps, S., & Kriner, D. (2023). How AI Threatens Democracy. *Journal of Democracy*, 34(4), 122–131. <https://doi.org/10.1353/jod.2023.a907693>
- Lainjo, B. (2023). THE GLOBAL SOCIAL DYNAMICS AND INEQUALITIES OF ARTIFICIAL INTELLIGENCE. *International Journal of Innovation Scientific Research and Review*, 05, 4966–4974. <http://www.journalijisr.com>
- Limantè, A. (2024). Bias in Facial Recognition Technologies Used by Law Enforcement: Understanding the Causes and Searching for a Way Out. *Nordic Journal of Human Rights*, 42(2), 115–134. <https://doi.org/10.1080/18918131.2023.2277581>
- Lindman, J., Makinen, J., & Kasanen, E. (2023). Big Tech's power, political corporate social responsibility and regulation. *Journal of Information Technology*, 38(2), 144–159. <https://doi.org/10.1177/02683962221113596>
- Longpre, S., Mahari, R., Lee, A., Lund, C., Oderinwale, H., Brannon, W., Saxena, N., Obeng-Marnu, N., South, T., Hunter, C., Klyman, K., Klamm, C., Schoelkopf, H., Singh, N., Cherep, M., Anis, A. M., Dinh, A., Chitongo, C., Yin, D., ... Pentland, S. (2024). Consent in Crisis: The Rapid Decline of the AI Data Commons. 38th Conference on Neural Information Processing Systems. <https://doi.org/10.48550/arXiv.2407.14933>
- Maphosa, V. (2024). The Rise of Artificial Intelligence and Emerging Ethical and Social Concerns. *AI, Computer Science and Robotics Technology*, 3. <https://doi.org/10.5772/acrt.20240020>
- Mcmullen, T. (2020). Unconscious Bias on the Implementation and Utilization of Emerging Technologies by Law Enforcement Agencies, and Effects on the Security and Privacy of Citizens in Florida: A Case Study of Florida.
- Mitra, S., Behera, I., Arun, A., & Patnaik, S. (2025). Generative Artificial Intelligence Vis-À-Vis: Social, Ethical and Legal Issues (pp. 121–146). https://doi.org/10.1007/978-3-031-87252-5_10
- Neroni Rezende, I. (2022). Facial Recognition for Preventive Purposes: The Human Rights Implications of Detecting Emotions in Public Spaces (pp. 67–98). https://doi.org/10.1007/978-3-031-13952-9_4
- Njoto, S., Cheong, M., Lederman, R., McLoughney, A., Ruppanner, L., & Wirth, A. (2022). Gender Bias in AI Recruitment Systems: A Sociological- and Data Science-based Case Study. 2022 IEEE International Symposium on Technology and Society (ISTAS), 1–7. <https://doi.org/10.1109/ISTAS55053.2022.10227106>
- Owolabi, O. O., Adewusi, O. B., Ajayi, F. A., Asunmonu, A. A., Chukwurmazu, O., Ederhion, J., & Ayeni, O. M. (2025). Developing Frameworks for Assessing and Mitigating Bias in AI Systems: A Case Study on Ensuring Fairness in AI Diagnostic Tools through Diverse Training Datasets to Prevent Misdiagnosis in Underrepresented Populations. *South Asian Research Journal of Engineering and Technology*, 7(01), 33–38. <https://doi.org/10.36346/sarjet.2025.v07i01.004>
- Pöhler, L., Schrader, V., Ladwein, A., & von Keller, F. (2024). A Technological Perspective on Misuse of Available AI. <http://arxiv.org/abs/2403.15325>
- Prasanth, A., John Vadakkan, D., Surendran, P., & Thomas, B. (2023). Role of Artificial Intelligence and Business Decision Making. *IJACSA International Journal of Advanced Computer Science and Applications*, 14(6), 965–969. www.ijacsa.thesai.org
- Prince, N. U., Faheem, M. A., Khan, O. U., Hossain, K., Uddin Prince, N., Alkhayyat, A., Hamdache, A., & Elmouki, I. (2024). AI-Powered Data-Driven Cybersecurity Techniques: Boosting Threat Identification and Reaction. *Nanotechnology Perceptions*, 20(10), 332–353. <https://doi.org/10.13140/RG.2.2.22975.52644>
- Pulivarthy, P., & Whig, P. (2024). Bias and Fairness Addressing Discrimination in AI Systems. In *Ethical Dimensions of AI Development* (pp. 103–126). <https://doi.org/10.4018/979-8-3693-4147-6.ch005>
- Rainy, T. A. (2025). *International Journal of Scientific Interdisciplinary Research*. *International Journal of Scientific Interdisciplinary Research*, 06(01), 28–59. <https://doi.org/10.63125/0k4k5585>

- Ryan, M., & Stahl, B. C. (2021). Artificial intelligence ethics guidelines for developers and users: clarifying their content and normative implications. *Journal of Information, Communication and Ethics in Society*, 19(1), 61–86. <https://doi.org/10.1108/JICES-12-2019-0138>
- Selbst, A. D., Balkin, J., Belt, R., Brennan-Marquez, K., Calo, R., Crotoft, R., Davison, P., Emerson, B., Ferryman, K., Grady, M., Grimmelmann, J., Horwitz, J., Katyal, S., Latonero, M., Mulligan, C., Nguyen, A., Parson, T., Patel, S., Price, N., ... Wexler, R. (2020). NEGLIGENCE AND AI'S HUMAN USERS. In *BOSTON UNIVERSITY LAW REVIEW* (Vol. 100). <http://creativecommons.org/licenses/by-sa/4.0/>.The
- Seyyed-Kalantari, L., Zhang, H., McDermott, M. B. A., Chen, I. Y., & Ghassemi, M. (2021). Underdiagnosis bias of artificial intelligence algorithms applied to chest radiographs in under-served patient populations. *Nature Medicine*, 27(12), 2176–2182. <https://doi.org/10.1038/s41591-021-01595-0>
- Shalevska, E., & Walker, A. (2025). Are AI Models Politically Neutral? Investigating (Potential) AI Bias Against Conservatives. In *International Journal of Research Publication and Reviews Journal homepage: www.ijrpr.com* (Vol. 6). www.ijrpr.com
- Sholademi, D. B. (2024). Leveraging AI for Detecting Deep Fakes and Combating Financial Fraudulent Identity Schemes. *International Journal of Research Publication and Reviews*, 5(12), 4096–4111. <https://doi.org/10.55248/gengpi.5.1224.250131>
- Singh, T. (2024). AI-Driven Surveillance Technologies and Human Rights: Balancing Security and Privacy (pp. 703–717). https://doi.org/10.1007/978-981-97-3690-4_53
- Smuha, N. A. (2021). From a 'race to AI' to a 'race to AI regulation': regulatory competition for artificial intelligence. *Law, Innovation and Technology*, 13(1), 57–84. <https://doi.org/10.1080/17579961.2021.1898300>
- Solarova, S., Podroužek, J., Mesarčík, M., Gavorník, A., & Bielikova, M. (2023). Reconsidering the regulation of facial recognition in public spaces. *AI and Ethics*, 3(2), 625–635. <https://doi.org/10.1007/s43681-022-00194-0>
- Sophia LI. (2025). The Social Harms of AI-Generated Fake News: Addressing Deepfake and AI Political Manipulation. *Digital Society & Virtual Governance*, 1(1), 72–88. <https://doi.org/10.6914/dsvg.010105>
- Špelda, P. (2024). Optimizing Hybrid Decision-Making Models in AI-Integrated Weapon Systems: Balancing Human Control, Ethical Oversight, and Efficiency through AI Autonomy.
- Stamova, I., & Draganov, M. (2020). Artificial Intelligence in the Digital Age. *IOP Conference Series: Materials Science and Engineering*, 940, 012067. <https://doi.org/10.1088/1757-899X/940/1/012067>
- Stone, M., Aravopoulou, E., Ekinci, Y., Evans, G., Hobbs, M., Labib, A., Laughlin, P., Machtynger, J., & Machtynger, L. (2020). Artificial intelligence (AI) in strategic marketing decision-making: a research agenda. *The Bottom Line*, 33(2), 183–200. <https://doi.org/10.1108/BL-03-2020-0022>
- Taiwo, A. (2025). AI-Driven Fact-Checking in Journalism: Enhancing Information Veracity and Combating Misinformation: A Systematic Review. *SSRN*.
- Treleaven, P., Barnett, J., Brown, D., Bud, A., Fenoglio, E., Kerrigan, C., Koshiyama, A., Sfeir-Tait, S., & Schoernig, M. (2023). The Future of Cybercrime: AI and Emerging Technologies Are Creating a Cybercrime Tsunami. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.4507244>
- Veerasamy, N., & Pieterse, H. (2022). Rising Above Misinformation and Deepfakes. *International Conference on Cyber Warfare and Security*, 17(1), 340–348. <https://doi.org/10.34190/iccws.17.1.25>
- Yen, C.-P., & Hung, T.-W. (2021). Achieving Equity with Predictive Policing Algorithms: A Social Safety Net Perspective. *Science and Engineering Ethics*, 27(3), 36. <https://doi.org/10.1007/s11948-021-00312-x>
- Zajko, M. (2022). Artificial intelligence, algorithms, and social inequality: Sociological contributions to contemporary debates. *Sociology Compass*, 16(3). <https://doi.org/10.1111/soc4.12962>
- Zeng, J. (2020). Artificial intelligence and China's authoritarian governance. *International Affairs*, 96(6), 1441–1459. <https://doi.org/10.1093/ia/iaa172>