

The Migration of Business Ethics within Artificial Intelligence (AI)

Khaled M. Musa
College of Business
Sacramento State University
Sacramento, California
K.musa@csus.edu

Cecilia Walsh
College of Business
Sacramento State University
Sacramento, California
ceciliarwalsh@csus.edu

Laila Akhavi
College of Business
Sacramento State University
Sacramento, California
lailaakhavi@csus.edu

Nasirullah Rekhteen
College of Business
Sacramento State University
Sacramento, California
nasirullahrekhteen@csus.edu

Abstract— Artificial intelligence (AI) transforms business operations across industries, offering unprecedented opportunities for efficiency, innovation, and decision-making. However, its rapid deployment also raises critical ethical concerns, including algorithmic bias, privacy violations, and accountability gaps. This paper examines the ethical dimensions of AI integration in business, drawing on historical precedents, current regulatory frameworks—such as the European Union’s AI Act and the U.S. Blueprint for an AI Bill of Rights—and corporate case studies, including IBM’s AI-powered hiring tools and Apple’s credit-scoring algorithm. The analysis reveals a consistent pattern: while AI holds transformative potential, inadequately regulated systems risk perpetuating discrimination, reinforcing structural inequities, and eroding public trust. Central to addressing these risks are the ethical principles of transparency, fairness, and human oversight. The study explores the tension between technological innovation and ethical responsibility, particularly in the context of biased data and algorithms that can disproportionately impact marginalized communities. It also examines emerging technologies such as generative AI and natural language processing (NLP), emphasizing their dual capacity for social progress and ethical harm. Recommendations include implementing proactive governance strategies, such as multidisciplinary ethics boards, robust bias-detection mechanisms, and continuous stakeholder engagement. These measures can help organizations align AI development with broader societal values and regulatory expectations. The paper concludes that ethical AI is not merely a compliance obligation but a strategic imperative. By embedding ethical considerations into the design, deployment, and oversight of AI systems, businesses can foster trust, ensure fairness, and achieve sustainable success in the digital era..

Keywords—business Ethics, Artificial Intelligence, Business and AI, Ethical AI Standards

Artificial intelligence (AI) is transforming modern business practices by enhancing operational efficiency, improving decision-making processes, and personalizing customer experiences. However, the rapid and widespread adoption of AI has introduced significant ethical challenges, including algorithmic bias, privacy infringements, and accountability deficits. While AI-driven innovations such as predictive analytics and automation offer substantial competitive advantages, high-profile cases—such as discriminatory hiring algorithms and the misuse of facial recognition technologies—underscore the urgent need for ethical oversight and responsible implementation [7].

The growing reliance on AI across business functions makes the ethical discussion increasingly urgent. As AI systems shape critical areas such as hiring, marketing, finance, and customer engagement, their societal impact is both far-reaching and consequential. Recent regulatory developments, including the European Union’s AI Act (2024) and the U.S. Blueprint for an AI Bill of Rights (2022), signal a global recognition of the need for ethical governance. Businesses that proactively address these concerns not only reduce legal and reputational risks but also foster public trust and ensure long-term viability. This paper argues that ethical AI is not merely a matter of regulatory compliance, but a strategic imperative in the digital age.

The purpose of this research is to examine the ethical challenges posed by the usage of artificial intelligence in business and to propose strategies for balancing technological advancement with moral responsibility. Through an analysis of existing ethical frameworks, corporate case studies, and evolving regulatory approaches, this paper seeks to offer actionable insights for policymakers, developers, and business leaders. Ultimately, it highlights the critical importance of embedding core ethical principles—such as transparency, fairness, and accountability—into the design, development, and deployment of AI systems.

I. INTRODUCTION

II. HISTORY OF ETHICS IN AI

1. Background and History of Ethics in AI

The ethical implications of artificial intelligence (AI) have evolved in parallel with its technological development, rooted in longstanding philosophical debates about autonomy, responsibility, and the role of machines in society. Early pioneers, such as Alan Turing, questioned whether machines could truly "think," sparking foundational discussions about machine intelligence and moral agency. Isaac Asimov's Three Laws of Robotics further contributed to the ethical discourse by introducing conceptual safeguards for autonomous systems.

However, it was not until the 21st century—driven by the emergence of big data, machine learning, and algorithmic decision-making—that these ethical concerns became pressing. Real-world incidents have underscored the potential harm of poorly governed AI. For example, Microsoft's Tay chatbot, which quickly adopted offensive language and behaviors, and biased AI recruitment tools that discriminated against women, revealed how AI can perpetuate societal biases when deployed without oversight. These failures helped catalyze a global reckoning over AI ethics, emphasizing the need for accountability, transparency, and human oversight in AI development and deployment [21].

2. Current Ethical Business Issues in AI

In response to the growing influence of AI, governments, corporations, and research institutions are developing frameworks to promote ethical and responsible deployment. The European Union's AI Act (2024) classifies AI systems based on risk levels—ranging from minimal to unacceptable—and imposes rigorous transparency and accountability requirements for high-risk applications. In the United States, the Blueprint for an AI Bill of Rights (2022) outlines five key principles, including protection from algorithmic discrimination, data privacy, and human alternatives, to safeguard the public interest [23].

Corporate actors are also taking initiative. Companies such as Google and IBM have established internal AI ethics boards and created fairness toolkits to audit algorithms for bias, ensuring alignment with internal values and public expectations. On a global scale, UNESCO's Recommendation on the Ethics of Artificial Intelligence (2021) advocates for international cooperation to uphold human rights, inclusiveness, and environmental sustainability in AI development.

Collectively, these efforts signal a growing consensus: ethical AI cannot be achieved through reactive measures alone. Instead, it requires proactive, transparent, and collaborative governance to ensure AI technologies serve the broader public good.

3. General Ethical Guidelines for AI

A number of foundational principles underpin contemporary ethical AI frameworks. These guidelines aim to ensure that AI systems are not only technically robust but also aligned with societal values:

- **Transparency:** AI systems should be explainable and understandable to users and stakeholders. Techniques such as explainable AI (XAI) are increasingly employed to make algorithmic decisions interpretable.
- **Fairness:** Efforts must be made to identify and mitigate biases in both data and algorithmic outcomes. Tools like IBM's AI Fairness 360 Toolkit are designed to detect and reduce discriminatory patterns in machine learning models.
- **Accountability:** There must be clear mechanisms to assign responsibility and liability for AI-driven errors. Regulations such as the EU's General Data Protection Regulation (GDPR) support this through provisions like the "right to explanation" [10].
- **Privacy:** AI systems must comply with data protection laws, ensuring the secure and ethical use of personal information. This includes practices like data anonymization during model development.
- **Human Oversight:** Humans should retain meaningful control over AI systems, particularly in high-stakes applications such as healthcare, criminal justice, and finance [11].

These principles are codified in global frameworks such as the OECD AI Principles and the IEEE's Ethically Aligned Design, which provide high-level guidance for ethical AI development. However, implementation remains uneven across sectors, often depending on organizational priorities, regulatory environments, and technological maturity.

4. Ethical Challenges of AI in Business

Businesses face a range of complex ethical dilemmas in the deployment of artificial intelligence, often due to the intersection of profit-driven objectives and socially impactful technologies:

- **Bias in Hiring and Lending:** AI-powered recruitment tools and loan approval algorithms have been shown to replicate or amplify existing societal biases, leading to discriminatory outcomes for marginalized groups.
- **Surveillance Capitalism:** The use of AI to track consumer behavior, preferences, and personal data—often without informed consent—raises significant

privacy concerns and contributes to exploitative data practices [24].

- **Job Displacement:** Automation and AI-driven efficiency gains can lead to workforce reductions, disproportionately affecting low-skilled or routine job sectors and exacerbating socioeconomic inequality.
- Many AI systems, particularly those based on deep learning, operate with limited explainability. This opacity can be especially problematic in high-stakes domains like credit scoring or healthcare, where stakeholders are unable to understand or contest decisions.

Addressing these challenges requires a multidisciplinary approach that combines technical interventions—such as bias audits and explainability tools—with broader policy reforms and active stakeholder engagement. Only through collaboration between technologists, ethicists, regulators, and affected communities can businesses develop AI systems that are both innovative and ethical.

III. NEW TECHNOLOGY IN AI

1. The Rise of Generative AI and Its Ethical Implications

Generative AI represents one of the most rapidly advancing branches of artificial intelligence, with applications spanning industries such as healthcare, architecture, marketing, and enterprise operations [4]. These models—based on deep learning—can produce original text, images, audio, and even predictive visualizations vast datasets [14]. A widely known example is ChatGPT, which can generate essays, poems, and song interpretations, demonstrating the power and flexibility of natural language generation.

In healthcare, generative AI is being used to enhance diagnostic imaging by interpreting X-rays, generating clinical reports, and simulating disease progression over time. These tools not only improve diagnostic accuracy but also streamline workflows for medical professionals. In marketing, generative AI is increasingly employed to develop personalized campaigns and promotional content, with projections indicating that it will contribute to the creation of approximately 30% of marketing materials by the end of 2025 [3].

However, generative AI is also reshaping the workforce. By automating tasks traditionally performed by humans, it contributes to job displacement in various sectors. According to McKinsey, industries that fail to integrate AI technologies—particularly in media, communications, and technology—risk becoming obsolete [20]. Notably, the same research suggests that employees may be more prepared for AI-driven changes

than many organizational leaders anticipate. Despite increasing investment in generative AI, only 1% of respondents believe their organizations have achieved a mature implementation level [15].

The expansion of generative AI also raises significant ethical concerns. Since these models use large datasets, they may reproduce or amplify harmful biases based on data. This can result in outputs that are offensive, discriminatory, or misleading. Additionally, many generative AI models depended on personal or proprietary data without proper consent, raising serious questions about privacy and data protection. To address these issues, it is essential that AI systems are developed with transparency, explainability, and privacy safeguards, ensuring outputs are trustworthy and ethically aligned.

2. The Growth of Natural Language Processing (NLP) and Associated Ethical Concerns

Another rapidly evolving AI technology is Natural Language Processing (NLP), which enables machines to understand, interpret, and generate human language. NLP powers a wide range of applications, including text summarization, language translation, and conversational AI. One of its most prominent uses is virtual chat assistants, where it is enhancing customer service by enabling more natural, personalized, and emotionally intelligent interactions. Through sentiment analysis, NLP systems can interpret the emotional tone of user input, allowing businesses to tailor responses and improve user satisfaction.

Voice-activated assistants like Amazon's Alexa and Apple's Siri are examples of NLP in action, enabling real-time, voice-based interactions. These technologies improve accessibility by allowing users—regardless of technical ability—to issue voice commands and receive spoken responses, simulating a natural conversation.

Despite its capabilities, NLP introduces a number of ethical challenges. Like other AI systems, NLP models are trained on large volumes of data, which can embed existing societal biases into the system. These biases may be reflected in how the model interprets or generates language, potentially reinforcing stereotypes or producing unfair outcomes. Privacy is another key concern, as NLP applications often process sensitive or confidential information. Improper handling of such data—especially in customer service or healthcare contexts—can result in serious breaches of trust and regulatory non-compliance [9].

Looking ahead, NLP is expected to evolve through advancements in semantic and cognitive technologies, with the goal of enabling machines to understand context, intent, and meaning more accurately. The ultimate aim is to create

intelligent, efficient, and user-friendly platforms that support smarter search, seamless communication, and more inclusive digital experiences [12].

3. The Growing Importance of AI Regulation and Governance

As artificial intelligence continues to evolve at a rapid pace, regulation has become a central focus for governments and regulatory bodies worldwide. Ensuring the ethical use of AI requires structured oversight, which is where AI governance teams play a vital role. These teams are responsible for implementing policies, managing risk, and ensuring compliance with emerging laws and ethical standards.

One landmark regulatory effort is the European Union's Artificial Intelligence Act, passed in 2024. Modeled in part after the General Data Protection Regulation (GDPR)—which came into effect in 2018 to give individuals greater control over their personal data—the AI Act classifies AI systems by risk level and enforces transparency and accountability standards. Like GDPR, the AI Act is expected to set a global precedent, encouraging other countries to adopt similar legislation.

Additionally, the international community has seen increased collaboration on AI standards and regulatory frameworks, with cross-jurisdictional efforts to align ethical principles and technical requirements. As policymakers gain a deeper understanding of AI's societal and economic risks, the momentum for robust regulatory frameworks is expected to grow [6].

This regulatory evolution also underscores the rising demand for AI professionals equipped to implement responsible AI systems. As organizations increase their AI investments—92% of companies plan to do so within the next three years [16]—there will be a corresponding need for individuals skilled in governance, compliance, and ethical AI design. Building this workforce will be critical to ensure that AI technologies are developed and deployed in ways that are both innovative and accountable.

IV. CASE STUDIES

1. Case Study: IBM's AI Hiring Assistant and Ethical Adaptation'

A notable case study highlighting how AI is being shaped by ethical considerations involves IBM, a multinational technology company specializing in cloud services, software, and hardware consulting. IBM developed an AI-powered hiring assistant designed to support various recruitment tasks, including generating job descriptions, sourcing candidates, and conducting candidate outreach. These functions traditionally fall under human

recruiters, but the AI tool aims to streamline and optimize the hiring process.

However, the use of AI in hiring raises significant ethical concerns—particularly around algorithmic bias. Critics have expressed concern that such tools may reinforce existing social inequalities unless carefully monitored and regulated. For instance, a study conducted by Kyra Wilson, a student at the University of Washington's Information School, investigated how AI-driven hiring tools might discriminate against certain occupations and demographic groups. Her research found that societal biases embedded in the data led to outcomes that disproportionately favored white and male candidates.

This type of research underscores the importance of transparency and proactive mitigation of bias in AI systems. According to the developers behind IBM's hiring assistant, one of the most effective strategies to address these issues lies in the large language model (LLM) training stage, where efforts can be made to reduce bias in the underlying data. However, they also acknowledge that some degree of bias may still persist, highlighting the need for continuous auditing and human oversight throughout the deployment process.

To address these concerns, regulatory measures are emerging. For example, New York State law now requires companies using automated hiring tools to involve human reviewers in the candidate evaluation process after AI systems generate shortlists. In parallel, companies are developing tools like Granite Guardian 3.0, designed to detect and flag biased content within AI outputs.

Additionally, the Biden-Harris Administration has taken steps toward establishing national ethical standards for AI, issuing an executive order with guidelines aimed at promoting fairness, transparency, and accountability in AI use [22]. These developments suggest a shift toward stricter oversight, reinforcing the need for organizations to integrate ethical considerations into AI systems from the outset.

2. Case Study: Algorithmic Bias in Apple's Credit Card

Despite efforts by many companies to embed ethical principles into their AI systems, there have been high-profile cases where algorithmic bias still emerged—highlighting the challenges of achieving truly fair AI. A notable example involves Apple Inc.'s credit card, launched in partnership with Goldman Sachs. Shortly after its release, the credit card faced public scrutiny over claims that it discriminated against users based on gender.

One widely cited incident involved an entrepreneur who reported on Twitter that he was granted a credit limit 20 times higher than his wife's, despite them sharing identical financial profiles—including credit scores and income. This raised concerns that the AI algorithm responsible for determining credit limits might be biased. Although Goldman Sachs denied the presence of gender bias, the

controversy prompted investigations by regulators and generated widespread media attention [13].

Importantly, gender was not explicitly included as an input variable in the algorithm. However, the model may have inferred gender indirectly from other correlated data, leading to biased outcomes. This case underscores a critical issue in AI ethics: even when protected characteristics are excluded, algorithms can still reflect societal biases embedded in training data or learned correlations.

The incident revealed the difficulty of building truly unbiased AI systems, especially in high-stakes industries like finance. It also illustrated the gap between algorithmic decision-making and societal expectations of fairness and transparency. As financial institutions increasingly turn to AI for credit scoring, loan approvals, and other key decisions, concerns are growing that regulatory frameworks may struggle to keep pace with technological advancements [17].

This case reinforces the need for strong oversight, bias auditing tools, and regulatory safeguards to minimize discriminatory outcomes in AI-driven decision-making. It also emphasizes the importance of explainability and human oversight—particularly in industries where fairness is both a legal and ethical imperative.

V. NEW PROSPECTS

1. Challenges in Enforcing Ethical AI Practices

The rapid advancement of artificial intelligence has compelled governments, corporations, and institutions to consider the ethical implications of AI design and deployment. However, implementing these ethical principles presents a complex set of challenges. Organizations often struggle to operate abstract values such as fairness, transparency, accountability, and consistency into tangible policies and workflows.

One significant barrier is the inconsistency of regulations across jurisdictions, industries, and use cases, which complicates compliance and hinders effective oversight. As one study notes, “the inconsistency of regulations across multiple jurisdictions, industries, and use cases poses challenges not just for compliance, but also for enforcement by regulators” [3]. This fragmented regulatory environment results in conflicting rules and expectations, making it difficult for companies to adhere to a uniform ethical standard.

Furthermore, many AI regulations lack robust enforcement mechanisms, instead relying on self-regulation by organizations. As another source highlights, this reliance on internal governance “carries intrinsic risks of conflicts of interest” [5]. Without independent external oversight, companies may prioritize business interests—such as profitability or speed to market—over their ethical

obligations. This self-policing model risks undermining the credibility and effectiveness of ethical AI initiatives.

Ultimately, the absence of strong, harmonized regulation underscores the need for coordinated policy efforts and independent monitoring structures that can hold organizations accountable and ensure that ethical commitments are upheld in practice.

2. From Traditional Governance to AI Ethics Frameworks

Prior to the rise of artificial intelligence, companies relied on traditional corporate governance mechanisms to uphold ethical standards. These included codes of conduct, compliance officers, corporate social responsibility (CSR) programs, and whistleblower systems—all designed to promote accountability, fairness, and transparency in business practices. These tools were often reinforced through employee training and compliance with industry-specific regulations.

However, as AI has become increasingly embedded in core business functions—particularly in decision-making processes—traditional governance structures have proven insufficient to address emerging challenges. Issues such as algorithmic bias, lack of explainability, and automated decision-making without human oversight demand new approaches to ethical risk management.

In response, many organizations are now adopting AI-specific ethics governance frameworks that build upon traditional structures while introducing specialized roles and policies. These may include AI ethics boards, designated ethics leads within business units, and employee advocacy networks that foster responsible AI practices across the organization. These frameworks aim to ensure that foundational ethical principles—particularly fairness, transparency, and accountability—remain central in an increasingly automated and data-driven business environment.

By evolving their governance systems, forward-looking companies are better positioned to navigate the ethical complexities of AI while maintaining public trust and regulatory compliance.

3. The Challenge of Scale and Emerging Ethical Risks

One of the most pressing challenges in ethical AI implementation is addressing concerns that emerge over time—especially as AI systems scale rapidly across industries. As organizations adopt AI in diverse applications, from recruitment to finance to surveillance, the complexity of these systems makes it increasingly difficult to detect and prevent ethical failures before they occur. Given the speed of development and deployment, even well-governed AI systems are likely to overlook critical issues.

A particular concern is how AI models can inadvertently learn and reinforce subtle human biases. For example, generative AI and other advanced models may pick up on nuanced cues—such as facial expressions in loan decisions, personality traits in hiring, or behavioral patterns in online proctoring—that can lead to discriminatory outcomes. These biases often go unnoticed until they affect real users, particularly marginalized or underrepresented groups, highlighting the importance of ongoing monitoring and ethical review.

This challenge is further complicated by the fact that many ethical risks do not emerge until after deployment, making them difficult to anticipate through standard testing or oversight. In such cases, regulatory standards and governance mechanisms often lag behind technological innovation, allowing problematic systems to persist before corrective action is taken.

While AI technologies offer significant advantages—such as efficiency, scalability, and personalization—they also introduce legal, ethical, and reputational risks that can undermine public trust and organizational integrity. Proactively addressing these challenges requires not only robust design and governance but also a commitment to continuous oversight and ethical adaptation as new use cases and risks emerge.

4. Legal, Ethical, and Reputational Risks of AI

While AI systems offer considerable advantages—such as automation, improved efficiency, and data-driven decision-making—they also pose significant legal, ethical, and reputational risks. Legally, organizations must navigate a rapidly evolving landscape of AI regulations. For instance, the European Union's Artificial Intelligence Act imposes strict oversight on high-risk AI systems, while in the United States, several states are introducing their own ethical AI legislation to protect individuals and ensure accountability.

A key example of legal and ethical risk is seen in the case of Amazon's AI-powered hiring tool, which was eventually abandoned. According to a report from CIO, the tool was designed to pre-screen candidates but ultimately "unintentionally favored male candidates over female ones due to biases in the training data." As a result, qualified female applicants were systematically disadvantaged, leading to indirect discrimination and public criticism [16]. This highlights how biased data can lead to real-world harm and legal liabilities.

Reputational risks also play a critical role in shaping public trust in AI. Misuse of AI technologies—such as deepfakes and unauthorized impersonations—can damage brand reputation and erode customer loyalty. In the United States, the House of Representatives recently passed the "Take It Down" Act, which criminalizes non-consensual deepfake pornography and mandates that platforms remove

such content. The law is a response to the surge in AI-generated illicit imagery, which has affected celebrities, public figures, and private individuals alike [4].

While not all uses of generative AI are harmful, incidents like these illustrate the urgent need for regulation, public awareness, and organizational safeguards. To mitigate legal and reputational damage, businesses must adopt ethical frameworks and establish robust AI governance programs that ensure fairness, transparency, and oversight throughout the AI lifecycle. By proactively managing these risks, companies can harness AI's benefits while maintaining ethical integrity and public trust.

VI. RECOMMENDED IMPLICATIONS

Given the rapid integration of artificial intelligence into business operations, it is imperative that organizations adopt proactive and comprehensive ethical governance frameworks to manage emerging risks and uphold public trust. As outlined in an IBM report, an effective AI governance structure consists of four core roles [18]:

Policy Advisory Committee – Sets the strategic direction for AI ethics, privacy, and regulatory compliance at a global level.

AI Ethics Board – Comprised of senior ethics and privacy officers, this board defines, advises on, and enforces company-wide AI ethics policies.

AI Ethics Focal Points – Positioned within individual business units, these representatives serve as the first line of ethical oversight, identifying risks and ensuring that AI use aligns with company values.

Advocacy Network – A grassroots group of employees who foster a culture of responsible AI and champion ethical practices across the organization.

These roles emphasize that effective AI governance extends beyond technical fixes; it must involve organizational culture, cross-functional collaboration, and leadership accountability. As demonstrated in cases involving biased hiring tools and discriminatory credit algorithms, ethical risks can have significant legal, social, and reputational consequences. Addressing these issues requires integrating ethical principles—such as transparency, fairness, and accountability—into the AI lifecycle from the outset.

In conclusion, businesses that lead with ethical practices will not only comply more effectively with evolving regulations but will also earn greater trust from customers, employees, and stakeholders. By embedding moral responsibility into innovation, organizations can navigate the complexities of AI with confidence and contribute to a more equitable and sustainable digital future.

VII. CONCLUSION

While artificial intelligence holds tremendous promises for enhancing business operations, it simultaneously introduces significant ethical, legal, and reputational risks that must not be overlooked. The history and evolution of AI ethics underscore the critical need for oversight, transparency, and fairness in AI development and deployment. As AI adoption expands, challenges such as algorithmic bias, privacy violations, and accountability gaps grow increasingly complex and consequential.

Case studies—such as IBM’s AI hiring assistant and Apple’s credit card algorithm—demonstrate that without proactive and deliberate measures, AI systems risk perpetuating harmful biases rather than mitigating them. Furthermore, the inconsistent regulatory landscape across jurisdictions complicates efforts to implement ethical AI practices effectively.

As AI technologies, including generative AI, continue to evolve rapidly, organizations must adopt comprehensive ethical governance frameworks that embed core principles of transparency, fairness, and accountability from the outset. Such frameworks will not only enable businesses to comply with evolving legal requirements but also foster public trust, safeguard human rights, and ensure that AI serves society responsibly in the years to come.

REFERENCES

- [1] Brynjolfsson, E., & McAfee, A. (2014). *The Second Machine Age*. W.W. Norton.
- [2] Burrell, J. (2016). How the machine ‘thinks’. *Big Data & Society*, 3(1).
- [3] Chen, J., Storch, V., & Kurshan, E. (2021). Beyond fairness metrics: Roadblocks and challenges for ethical AI in practice. *arXiv*. <https://arxiv.org/abs/2109.13961>
- [4] Chow, A. R. (2025, April 29). Congress just passed its first bill tackling AI harms. *Time*. <https://time.com/7277746/ai-deepfakes-take-it-down-act-2025/>
- [5] Coursera. (2024, March 15). 20 examples of generative AI applications across industries. <https://www.coursera.org/articles/generative-ai-applications>
- [6] Davison, A. (2024, December 10). How AI can help the recruitment process—without hurting it. IBM. <https://www.ibm.com/think/news/ai-in-recruiting?>
- [7] Dastin, J. (2018). Amazon scraps secret AI recruiting tool that showed bias against women. *Reuters*.
- [8] Domin, H. (2024, September 5). AI governance trends: How regulation, collaboration, and skills demand are shaping the industry. *World Economic Forum*. <https://www.weforum.org/stories/2024/09/ai-governance-trends-to-watch/>
- [9] Duzduran, A. (2023, August 12). Ethics of natural language processing. *Brass for Brain*. <https://medium.com/brass-for-brain/ethics-of-natural-language-processing-c2470c08c8a4>
- [10] European Parliament. (2018). *General Data Protection Regulation (GDPR)*.
- [11] Floridi, L., et al. (2018). *AI4People—An ethical framework for AI*. *Science and Engineering Ethics*, 24(4).
- [12] Harazim, A. (2024, March 5). Advancements in natural language processing (NLP). *Iteo*. <https://iteo.com/blog/post/advancements-in-natural-language-processing-nlp/>
- [13] IBM. (2024, April 10). What should an AI ethics governance framework look like? <https://www.ibm.com/think/insights/ai-governance-framework-ethics>
- [14] Martineau, K. (2023, April 20). What is generative AI? IBM Research Blog. <https://research.ibm.com/blog/what-is-generative-AI>
- [15] Mayer, H., Yee, L., Chui, M., & Roberts, R. (2025). Superagency in the workplace: Empowering people to unlock AI’s full potential. McKinsey & Company. <https://www.mckinsey.com/capabilities/mckinsey-digital/our-insights/superagency-in-the-workplace-empowering-people-to-unlock-ais-full-potential-at-work>
- [16] McKinsey & Company. (2024, April 2). What is generative AI? <https://www.mckinsey.com/featured-insights/mckinsey-explainers/what-is-generative-ai>
- [17] Technologies, D. (2021, March 4). How gender bias led to the scrutiny of the Apple Card. *Datatron*. <https://datatron.com/how-gender-bias-led-to-the-scrutiny-of-the-apple-card/>
- [18] The Apple Card algo issue: What you need to know about A.I. in everyday life. (2019, November 15). *CNBC*. <https://www.cnn.com/2019/11/14/apple-card-algo-affair-and-the-future-of-ai-in-your-everyday-life.html>
- [19] Thomas, M. (2025, January 28). The future of artificial intelligence. *Built In*. <https://builtin.com/artificial-intelligence/artificial-intelligence-future>
- [20] Treatment Episode Data Set: Discharges (TEDS-D): Concatenated, 2006 to 2009. (2013, August). U.S. Department of Health and Human Services, Substance Abuse and Mental Health Services Administration, Office of Applied Studies. <https://doi.org/10.3886/ICPSR30122.v2>
- [21] Turing, A. (1950). Computing machinery and intelligence. *Mind*, 59(236).
- [22] Wade, C. (2024, June 17). Navigating the ethical and legal risks of AI implementation. *CIO*. <https://www.cio.com/article/2149672/navigating-the-ethical-and-legal-risks-of-ai-implementation.html>
- [23] White House. (2022). *Blueprint for an AI Bill of Rights*.
- [24] Wolf, M. J., et al. (2017). *Ethics of AI and Robotics*. *Stanford Encyclopedia of Philosophy*.
- [25] Zuboff, S. (2019). *The age of surveillance capitalism*. *PublicAffairs*.