

Bounded Rationality and AI in Decision Making

– Krishnendu Das*

Assistant Professor & HoD. Economics, Lalbaba College, University of Calcutta, Kolkata

 krishnendudas20@gmail.com  <https://orcid.org/0000-0001-5039-7286>

– Anamika Choudhary

Professor, Dept. of Economics, Dr. Shakuntala Misra National Rehabilitation University, Lucknow

 achaudhary@dsrnru.ac.in  <https://orcid.org/0000-0001-9161-0784>



ARTICLE HISTORY

Paper Nomenclature: Review of Literature (RoL)

Paper Code: GJEISV18I1JM2026ROL3

Submission at Portal (www.gjeis.com): 02-Jan-2026

Manuscript Acknowledged: 12-Jan-2026

Originality Check: 27-Jan-2026

Originality Test (Plag) Ratio (DrillBit): 01%

Author Revert with Rectified Copy: 04-Feb-2026

Peer Reviewers Comment (Open): 06-Feb-2026

Single Blind Reviewers Explanation: 20-Feb-2026

Double Blind Reviewers Interpretation: 25-Feb-2026

Triple Blind Reviewers Annotations: 11-March-2026

Author Update (w.r.t. correction, suggestion & observation): 16-March-2026

Camera-Ready-Copy: 20-March-2026

Editorial Board Excerpt & Citation: 29-March-2026

Published Online First: 31-March-2026

ABSTRACT

Purpose: This paper explores the intricate relationship between human decision-making and artificial intelligence, viewed through the lens of Herbert Simon's bounded rationality theory. It investigates how AI simultaneously augments and constrains human cognition, while raising urgent ethical, environmental, and governance concerns that remain inadequately addressed in existing literature i.e. whether AI genuinely frees human cognition from its inherent limitations, or merely substitutes one set of constraints with another while generating fresh ethical and governance challenges in the process.

Design/Methodology/Approach: The study builds on a systematic review of over 300 sources spanning behavioural economics, cognitive psychology, AI ethics, neuroscience, and organizational theory including peer-reviewed journals, government reports, and industry white papers. The review follows the arc from Simon's original challenge to the myth of homo economicus, through Kahneman and Tversky's work on heuristics, and onward to AI's journey from early rule-based systems to today's deep learning models operating across healthcare, finance, and manufacturing.

Findings: AI does not overcome bounded rationality, it relocates it, shifting constraints from the human mind to computational systems. In complex domains, AI sharpens decision quality considerably. Yet it also breeds new forms of irrationality: entrenched algorithmic bias, quiet erosion of critical thinking, corporate-engineered choice environments, and platform designs built more around engagement than genuine user benefit. Where corporations hold the reins, these tendencies tend to deepen existing inequalities rather than correct them.

Originality: The paper puts forward the idea of AI-influenced rationality — arguing that AI does not sit passively beside human judgment but actively moulds it. This perspective bridges a genuine gap in the literature, connecting Simon's foundational theory with the lived realities of modern AI in a way that existing frameworks have not.

Paper Type: Review of Literature.

KEYWORDS: Decision Making Frameworks | Bounded Rationality | Incomplete Information | Biases Paradigm
Environmental Costs | AI Influencing Rationality

*Corresponding Author (Krishnendu)

- Present Volume & Issue (Cycle): Volume 18 | Issue-1 | Jan-Mar 2026
- International Standard Serial Number:
Online ISSN: 0975-1432 | Print ISSN: 0975-153X
- DOI (Crossref, USA) <https://doi.org/10.18311/gjeis/2026>
- Bibliographic database: OCLC Number (WorldCat): 988732114
- Impact Factor: 3.57 (2019-2020) & 1.0 (2020-2021) [CiteFactor]
- Editor-in-Chief: Dr. Subodh Kesharwani
- Frequency: Quarterly

- Published Since: 2009
- Research database: EBSCO <https://www.ebsco.com>
- Review Pedagogy: Single Blind Review/ Double Blind Review/ Triple Blind Review/ Open Review
- Copyright: ©2026 GJEIS and it's heirs
- Publishers: Scholastic Seed Inc. and KARAM Society
- Place: New Delhi, India.
- Repository (figshare): 704442/13

GJEIS is an Open access journal which access article under the Creative Commons. This CC BY-NC-ND license (<http://creativecommons.org/licenses/by-nc-nd/4.0>) promotes access and re-use of scientific and scholarly research and publishing.



Introduction

The 21st century has witnessed the evolution of Artificial Intelligence (AI) from a theoretical concept to cornerstone of how decisions are made from healthcare to trading floors, from business to government, profoundly impacting human lives, economies and societies as a whole.

Even with the rise of AI, the human decision making is still limited by cognitive constraints as articulated in the theory of bounded rationality of Herbert Simon. This paper shall review the available literature to fill the gap; the idea is that human decisions are inherently limited by incomplete information, finite computational capacity, and time scarcity (Simon, 1955). While AI improves our decision making ability, it also pronounces our existing cognitive limitations. Synthesizing about 200 peer-reviewed articles, 58 books, 23 government reports, and 15 industry white papers, will enable to address a critical gap in the literature: there is in fact no single complete framework that could explain the dual role of AI as both helping in human decision making and also amplifying its inherent constraints.

There is a need for such a review. Firstly, the use of AI has outpaced the scholarly consensus on its implications. While studies like Obermeyer et al. (2019) and Buolamwini & Gebru (2018) have illuminated algorithmic bias, few contextualize these findings within Simon's foundational theories. Second, the explosion of generative AI tools—ChatGPT, Gemini, and Claude—has reignited debates about machine autonomy, yet these discussions often neglect AI's historical role as a mirror of human decision-making strategies (Bostrom, 2014; Marcus, 2018). Third, policymakers lack actionable frameworks to balance AI's efficiency gains against risks like job displacement, privacy erosion, and cognitive deskilling (Manyika et al., 2017; Zuboff, 2019).

This paper bridges these gaps through five objectives:

- To trace the theoretical foundations of bounded rationality and their resonance in AI design.
- To understand the evolution of AI in reflecting and responding to human cognitive constraints.
- To examine the real world use of AI in various sectors.
- To address the ethical dilemmas created by AI, its impact on the environment as well as for long term risks.
- To suggest ways for hybrid human-AI which are fair, efficient and ecologically rational.

The review integrates ideas from varied fields viz. Behavioural economics, AI ethics and neuroscience. There is a biological parallel to AI's trade-offs between accuracy and computing limits (Hare et al., 2011; Griffiths et al., 2015). Reinforcement learning (an AI method) has a core conflict

between exploring new options and using what it already knows. This reflects a basic human behavior: the option of giving a trial to something novel out of curiosity or sticking with a safe chance to eliminate risk. (Blanchard et al., 2015; Sutton & Barto, 2018).

Structurally, the paper begins by traversing through bounded rationality's theoretical premise (Section 1), followed by an exploration of AI's decision-making paradigms (Section 2). Section 3 of the paper introduces the new idea of AI-influenced rationality. It explores how AI uses subtle prompts, or "algorithmic nudges," to change the way people think and make decisions. Sections 4 and 5 assimilates these theories in real-world applications and ethical challenges, while Section 6 projects equitable futures through hybrid systems and governance reforms. This review argues that AI isn't fundamentally different from human thinking; it's a digital reflection of it. By comparing classical and contemporary research, the paper is of immense help to academics, policymakers, and professionals to navigate through AI's benefits and risks with greater detail and ethical awareness.

Bounded Rationality

Herbert Simon's theory of bounded rationality, introduced in his seminal 1955 paper *A Behavioral Model of Rational Choice*, revolutionized the study of decision-making by challenging the neoclassical economic assumption of perfect rationality. Classical economics had conformed that humans act as *homo economicus*—maximising utility by utilizing all available information and becoming hyper-rational agents. Simon rejected this idealized view, and argued that human knowledge is constrained by three limitations: (1) **imperfect and incomplete information**, (2) **finite computational capacity**, and (3) **time scarcity** (Simon, 1955). Simon asserted that these constraints force decision-makers to adopt *satisficing* strategies meeting a minimum threshold of acceptability—rather opting for the ideal of optimization. This paradigm shift marked a transition from prescriptive models of behaving perfectly rational to descriptive accounts wherein people actually make decisions in the real world and reshape the fields.

A. Theoretical Foundations of Bounded Rationality

Simon's criticized that real-world decision-making deviates sharply from neoclassical ideals. In *Administrative Behavior* (1947), he highlighted that organizational decisions are seldom optimal and reflect pragmatic compromises under uncertainty. An example of satisficing is when managers take a shortcut in allocating money. Instead of calculating marginal returns, they adhere to what was given in the past or by departmental lobbying, thus saving both time and effort, and consider it as definitely a good solution (Simon, 1947). This work laid the groundwork for bounded rationality,



formalized in his 1955 paper, which introduced the concept of *procedural rationality*—the idea that decisions are shaped by the process of search and evaluation rather than predefined optimization rules.

Herbert Simon's framework centers on heuristics, or cognitive shortcuts, which allow us to make decisions despite limited information. For example, the **take-the-best heuristic** uses the most informative cue for rapid choices (Gigerenzer & Goldstein, 1996). However, these shortcuts can cause systematic biases. **Kahneman and Tversky's prospect theory** (1979) showed that people weigh potential losses more heavily than equivalent gains, demonstrating predictable deviations from pure rationality. **Heuristics, Biases, and Ecological Rationality**

The **heuristics and biases program** (Kahneman & Tversky, 1974) identified cognitive shortcuts that cause suboptimal decisions. For example, the **availability heuristic** leads people to overestimate events that are easily recalled (Tversky & Kahneman, 1973), while the **anchoring effect** shows how initial values influence judgments (Jacowitz & Kahneman, 1995).

In contrast, Gigerenzer's ecological rationality framework argues that heuristics are adaptive tools (Gigerenzer et al., 1999). The **recognition heuristic**, for instance, can outperform complex statistical models in certain contexts by leveraging familiarity (Goldstein & Gigerenzer, 2002). Similarly, the **fluency heuristic** explains why we favor easily processed options (Schwarz, 2004), reframing these shortcuts as evolved adaptations rather than cognitive flaws.

B. Organizational and Institutional Applications

Bounded rationality profoundly influenced organizational theory, where firms use standard operating procedures to simplify decision-making under cognitive limits (Cyert and March, 1963; Simon, 1972). Nudge theory applies this to policy, using “choice architectures” to guide decisions without restricting freedom (Thaler and Sunstein, 2008). In public policy, bounded rationality explains why policymakers rely on incremental adjustments (Lindblom, 1959). For example, climate change policies often prioritize politically feasible measures over optimal ones, reflecting “satisficing” in the face of complex constraints (Levin et al., 2012).

C. Computational and Mathematical Formalizations

Simon's ideas are formalized in computational models like Anderson's rational analysis (1990), which posits that cognitive systems optimize performance under constraints, mirroring computer cache hierarchies (Anderson & Schooler, 1991).¹ Griffiths et al. (2015) expanded this into resource-rational analysis, where decision-making is a cost-

benefit optimization balancing cognitive effort and accuracy (Vul et al., 2014), analogous to AI models like Monte Carlo tree search (Lewis et al., 2014). Critics argue this framework underestimates emotions; Damasio (1994) showed that emotional impairment hinders decisions, with emotions acting as guiding heuristics (Loewenstein et al., 2001). Social and cultural norms also shape decisions, as Henrich et al. (2001) found in collectivist societies. Furthermore, bounded rationality may overlook learning and adaptation, as individuals and organizations refine heuristics from experience (March, 1991), a process mirrored in machine learning (Steyvers et al., 2003).

D. AI and the Evolution of Bounded Rationality

Herbert Simon's concept of bounded rationality posits that humans use mental shortcuts. AI systems now transcend these cognitive limits, mediating complex decisions in fields like healthcare and finance. For instance, IBM Watson analyzes millions of medical studies to recommend treatments (Chen et al., 2021), and JPMorgan's COIN platform automates loan approvals (Davenport & Ronanki, 2018). These tools reduce cognitive effort, promising faster, more informed choices.

However, over-reliance on this “cognitive crutch” erodes critical thinking.¹ Studies show GPS navigation diminishes spatial memory (Ruginski et al., 2019), while AI tutors risk stifling creativity (Baker & Inventado, 2014).² AI also exploits cognitive biases to manipulate behavior. Ride-sharing apps use surge pricing (Chen & Sheldon, 2016), and e-commerce platforms use countdown timers (Kahneman et al., 1991), causing users to mistake algorithmic nudges for autonomous choices (Yeung, 2017). This psychological impact extends to social media, where reinforcement learning algorithms prioritize sensational content (Alter, 2017). Features like “infinite scroll” exploit psychological mechanisms linked to addiction (Harris, 2018), eroding self-regulation and fueling confirmation bias in echo chambers (Pariser, 2011).

E. Bounded Rationality in the Age of AI

The rise of AI has rekindled debates about rationality. Machine learning algorithms, such as deep neural networks, process vast datasets to identify patterns beyond human perception. However, these systems inherit human-like constraints. **Algorithmic satisficing** occurs when models settle for suboptimal solutions due to computational limits—training a neural network to 100% accuracy is often infeasible, so practitioners halt training once error rates plateau (Goodfellow et al., 2016). Reinforcement learning agents balance exploration (trying new strategies) and exploitation (maximizing known rewards), mirroring human trade-offs between curiosity and caution (Sutton & Barto, 2018).

Yet AI also faces **infrastructure bounds**. Training large models like GPT-4 requires immense energy and data,

creating environmental and ethical trade-offs (Strubell et al., 2019). Moreover, AI's reliance on historical data entrenches societal biases. Facial recognition systems trained on predominantly light-skinned datasets exhibit racial disparities, perpetuating discrimination in law enforcement (Buolamwini & Gebru, 2018). These limitations suggest that AI does not transcend bounded rationality but redistributes its constraints from cognition to computation.

F. Neuroscientific and Behavioral Insights

Neuroscience reveals satisficing activates the striatum, while optimizing engages the prefrontal cortex (Hare et al., 2011). This suggests heuristics conserve mental energy. Individuals dynamically adjust their strategies: high-stakes decisions prompt greater effort, while low-stakes ones rely on heuristics (Payne et al., 1993).

G. Synthesis and Future Directions

Bounded rationality remains a cornerstone of decision science, offering a pragmatic lens to understand adaptive intelligence. The theory's principles, like satisficing, still resonate, but it must evolve to address new constraints, such as algorithmic bias. Future research could explore hybrid human-AI systems that preserve human judgment or even quantum decision models (Biamonte et al., 2017). Ultimately, bounded rationality is a framework for understanding intelligence in a resource-constrained world. As Simon (1990) aptly noted, human capacity is small compared to complex problems, and acknowledging these bounds is crucial for navigating the digital age.

AI and Decision-Making

A. The Birth of Artificial Intelligence: Foundations of a Revolutionary Field

The conceptual origins of AI date back to the mid-20th century.¹ Alan Turing's 1950 paper introduced the Turing Test, a benchmark for machine intelligence, influenced by his wartime code-breaking work (Turing, 1950; Hodges, 1983).² The term "artificial intelligence" was coined at the 1956 Dartmouth Workshop, which set ambitious goals for machines to use language and reason abstractly (McCarthy et al., 1955).³ Early systems included Logic Theorist (Newell & Simon, 1956), which solved mathematical theorems, and ELIZA (Weizenbaum, 1966), a chatbot.⁴ These foundational efforts were supported by Cold War funding for military applications (Crevier, 1993).

B. The Evolution of AI: From Winters to the Deep Learning Revolution

AI's history is a story of cycles, swinging between periods of great hope and profound disappointment. The first major "AI winter" in the 1970s was triggered by critical reports, like the **ALPAC report** (1966) and the **Lighthill Report**

(1973), which deemed early AI efforts impractical and led to significant funding cuts (Hutchins, 1986; Lighthill, 1973).

The 1980s brought a resurgence with **expert systems**, which used rule-based frameworks to encode specific knowledge. Systems like MYCIN diagnosed bacterial infections with surprising accuracy (Shortliffe, 1976), and DENDRAL assisted in scientific discovery (Lindsay et al., 1980). However, these systems were "brittle," failing outside their narrow domains and struggling with the **knowledge acquisition bottleneck**, the difficulty of manually encoding expert knowledge (Feigenbaum, 1977).

A significant shift occurred in the 1990s toward **machine learning (ML)**, powered by new statistical models and greater computational power. Key innovations included Vladimir Vapnik's **support vector machines (SVMs)** (1995) for robust classification and Leo Breiman's **random forests** (2001), which improved accuracy by combining multiple decision trees (Cortes & Vapnik, 1995; Breiman, 2001). This era's highlight was IBM's **Deep Blue** (1997) defeating chess champion Garry Kasparov, a feat of brute-force computation that demonstrated the power of specialized, not general, AI (Campbell et al., 2002).

The **21st-century deep learning revolution** was fuelled by massive datasets and powerful hardware. Building on earlier work like **backpropagation** (Rumelhart et al., 1986), the 2012 breakthrough of **AlexNet** (Krizhevsky et al., 2012) marked a turning point. Later, **reinforcement learning** systems like DeepMind's **AlphaGo** (2016) defeated the world champion in Go (Silver et al., 2016). Recent generative models, such as **GPT-3** (2020), **Stable Diffusion** (2022), and **ChatGPT** (2023), have almost zeroed the line between human and machine creativity though ethical challenges still remain. (Brown et al., 2020; Bender et al., 2021; Patterson et al., 2021).

C. Human Decision-Making: From Rational Agents to Behavioral Realism

Daniel Kahneman and Amos Tversky's **prospect theory** (1979) razed classical rationalism, exhibiting that humans experience:

1. **Loss aversion:** (e.g., refusing to sell depreciating stocks to avoid realizing losses).
2. **Anchoring:** (e.g., estimating home prices based on listing figures).
3. **Framing effects:** Make inconsistent choices based on how options are presented.

Richard Thaler and Cass Sunstein's **nudge theory** (2008) applied these modalities to policymaking, propounding "libertarian paternalism" to usher decisions without any



force. For example, placing healthier foods at eye level in cafeterias nudges dietary choices without restricting freedom (Sunstein, 2014).

The Intersection of AI and Human Decision-Making

A. AI as a Mirror of Human Cognition

AI models imitate human cognition, like selective attention (Vaswani et al., 2017; Posner, 1980) and reinforcement learning (Gittins, 1979). However, they also acquire flaws. Black people are often misidentified by facial recognition. (Buolamwini & Gebru, 2018), and predictive policing targets minority neighborhoods, amplifying biases (Lum & Isaac, 2016).

B. AI as a Decision-Making Augmenter

Hybrid human-in-the-loop (HITL) systems synthesizes AI's computational power with human oversight. For example, IBM Watson for Oncology assists doctors, who retain final authority (Chen et al., 2021). Similarly, robo-advisors optimize portfolios while allowing user input (D'Acunto et al., 2019). However, a key risk is **automation bias**, where overtrust in AI can lead to dangerous outcomes, as exemplified by the Air France Flight 447 crash (Parasuraman & Manzey, 2010).

C. AI as a Disruptor of Human Agency

Autonomous systems, such as self-driving cars, face ethical dilemmas akin to the **trolley problem**—choosing between harming passengers or pedestrians (Awad et al., 2018). Unlike humans, who prioritize contextual empathy, AI relies on preprogrammed ethical frameworks (e.g., utilitarianism). The 2018 Uber autonomous vehicle fatality in Arizona, where a pedestrian was killed due to sensor failures, underscored accountability gaps in AI governance (NTSB, 2019).

D. AI-Influenced Rationality for Decision-Making: A Twist in the Story

Herbert Simon's theory of *bounded rationality*—the idea that humans make decisions under constraints of time, information, and cognitive capacity—remains foundational, but AI introduces both opportunities and ethical dilemmas. AI reshapes rationality by augmenting human decisions, challenging Herbert Simon's bounded rationality. However, its use by corporations risks subordinating individual autonomy to profit-driven interests. Equitable governance frameworks are essential to mitigate this harm and ensure AI serves the public good.

E. Corporate Influence and the Principal-Agent Problem

The principal-agent problem is magnified by corporate-controlled AI. Social media algorithms prioritize engagement over user well-being, as Meta executives ignored evidence of teen mental health crises to protect profits (Haig, 2021). Similarly, gig platforms manipulate drivers with opaque fare calculations and gamified incentives (Rosenblat & Stark, 2016). DoorDash's "stacked orders" feature, for instance, pressures drivers to accept unrealistic deliveries, highlighting how AI serves corporate profit at the expense of worker welfare (Ticona & Mateescu, 2018). Corporate data control entrenches power imbalances. Google and Meta dominate digital advertising by leveraging user data to prioritize paid content, while Amazon promotes private-label products over competitors (Epstein & Robertson, 2015; Khan, 2017). These practices create filter bubbles and enable misleading political microtargeting (Zuiderveen Borgesius et al., 2018).

Theoretical and Ethical Implications

A. Ecological Rationality and AI

Gerd Gigerenzer's **ecological rationality** (2008) posits that heuristics adapt to specific environments. AI systems could emulate this by tailoring strategies to contexts. For example:

- **Medical triage algorithms** prioritize patients based on urgency, mirroring emergency room protocols (Simon, 1972).
- **Financial "fast-and-frugal" trees** simplify credit scoring using minimal cues (Gigerenzer & Goldstein, 1996).

B. Explainability and Trust

The **black-box problem**—AI's opacity—undermines trust. **Explainable AI (XAI)** methods like LIME (Local Interpretable Model-agnostic Explanations) generate post hoc rationalizations for AI decisions, but these often lack causal fidelity (Arrieta et al., 2020). The EU's **General Data Protection Regulation (GDPR)** (2018) mandates a "right to explanation" for automated decisions, compelling industries to adopt interpretable models (Voigt & Von dem Bussche, 2017).

C. Ethical Governance and Sustainability

Global initiatives like the **OECD AI Principles** (2019) advocate for transparency and human-centric design. Regulatory sandboxes, such as Singapore's **AI Verify** (2022), allow for controlled testing (Cath et al., 2018). However, the environmental cost of AI is a concern; training GPT-3 emitted significant CO₂ (Patterson et al., 2021). The future lies in **human-AI symbiosis**, using machine precision to augment, not replace, human judgment, as Herbert Simon (1990) cautioned.

D. Ethical Implications and Structural Inequality

The corporate capture of AI exacerbates societal inequities.¹ **Surveillance capitalism** commodifies personal data for targeted advertising (Zuboff, 2019), and health insurers use fitness tracker data to adjust premiums (O'Neil, 2016), eroding trust. AI systems also encode historical biases; Amazon's recruitment tool was biased against women (Dastin, 2018), and predictive policing targets minority neighborhoods (Lum & Isaac, 2016). Furthermore, AI-generated deepfakes destabilize trust (Chesney & Citron, 2019). The environmental costs are substantial, with large models emitting significant CO₂ (Strubell et al., 2019), disproportionately burdening low-income communities with pollution (Hogan, 2021). These externalities highlight AI's unevenly distributed benefits.

E. Toward Equitable AI Governance

To address these challenges, systemic reforms are needed. Regulatory frameworks like the AI Act for European Union (2021) makes transparency mandatory for high-risk systems (European Commission, 2021). Features like Independent audits (Buolamwini & Gebru, 2018) and decentralized learning (Kairouz et al., 2021) grssroots movements (Benjamin, 2019) and ethical design standards (IEEE, 2019) are key. To promote equitable governance, public education (Raji et al., 2022) and labor unions challenging algorithmic management (Sundararajan, 2020) are equally critical.

F. Ethical Challenges of AI

With AI penetrating deeper into the society, the ethical challenges it brings are also serious which mandates due consideration and equitable governance.

i. Bias and Discrimination

AI systems can perpetuate biases in training data, decision making may lead to unfair outcomes in hiring, lending, and criminal justice (Harvard, 2020). Racial bias in facial recognition systems (Buolamwini and Gebru, 2018) and in health management algorithms (Obermeyer et al., 2019) reinforced the need for unbiased data.

ii. Lack of Transparency

The "black box" problem raises concerns about accountability and trust, highlighting the need for explainable AI (Capitol Technology University, 2023). The European Commission (2020) mandates the need for transparent AI systems to build trust (White Paper on Artificial Intelligence).

iii. Privacy and Data Protection

Relying on large datasets which includes sensitive personal information, requires robust data governance. Under the ethical Issues of Artificial Intelligence in Medicine

and Healthcare, Privacy and data protection challenges emphasizes informed consent. Floridi et al. (2019) propose an ethical framework for AI, addressing privacy risks.

iv. Economic Inequality and Job Displacement

Automation through AI risks job losses, particularly in routine tasks, potentially widening social inequalities, with WEF (2016) noting the widening wealth gap and the need for a post-labor economy (Top 9 Ethical Issues in Artificial Intelligence). Manyika et al. (2017) predict up to 800 million jobs could be lost to automation by 2030, necessitating reskilling (A Future That Works: Automation, Employment, and Labor).

v. Global Accessibility

Developing countries may lack access to advanced AI technologies, exacerbating disparities, with PMC (2022) highlighting accessibility issues in AI healthcare applications, noting challenges for low-income regions (Ethical Issues of Artificial Intelligence in Medicine and Healthcare). UNESCO (2021) recommends ethical AI to ensure global equity (Recommendation on the Ethics of Artificial Intelligence).

vi. Superintelligent AI Risks

The potential for artificial general intelligence (AGI) raises concerns about control and threats to humanity, with CoE (2022) identifying ethical concerns related to autonomous AI systems and responsibility apportionment (Common Ethical Challenges in AI). Bostrom (2014) discusses superintelligence risks, remaining relevant for current debates (Superintelligence: Paths, Dangers, Strategies).

Usmani, Happonen, and Watada (2022) propose a governance model for responsible AI, emphasizing fairness and explainability, while Machado, H., Silva, S., & Neiva, L. (2025) found public engagement in AI ethics aims to build trust, focusing on privacy and human control. Ricardo Baeza-Yates (2022), Li, F., Ruijs, N., & Lu, Y. (2023), and others highlight specific ethical issues like discrimination, transparency, and healthcare ethics, reinforcing the need for education and regulation.

AI and Its Real-Life Applications

Artificial intelligence (AI) has evolved from a theoretical concept to a practical tool transforming multiple sectors. Comparing over human judgement, it is free from emotional biases and makes objective decisions based on data and algorithms (Jatin Borana, 2016). Recent studies highlight AI's role across healthcare, finance, agriculture, retail, energy, and manufacturing, penetrating deeper into our daily lives (Rahmaniar, Maarif, Haq, et al., 2023).



A. Healthcare Sector

Bringing more accuracy in diagnosis, making treatment personalized and improving operational efficiency AI is revolutionizing healthcare. Medical images can be classified with better excellence which surpasses even experts in detecting conditions like pneumonia from X-rays or identifying cancer. Recent developments include **automated documentation** using NLP, streamlining clinical workflows (Smith et al., 2021). Predictive AI models analyze patient data to enable **personalized treatment plans**, improving outcomes (NIHR, 2023). In drug discovery, AI uses molecular simulations to predict toxicity, reducing costs and time to market (Brown et al., 2022). Overall, AI's transformative potential in medicine continues to grow (Hameem et al., 2019; Hussain and Qamar, 2016).

B. Finance Sector

In the financial sector, AI improves efficiency, accuracy, and security. It's used for **fraud detection** by analyzing real-time transaction data (Built In, 2023), and its identity verification platforms aid with **KYC compliance**. **AI-powered chatbots** and virtual assistants offer personalized financial advice, though regulatory scrutiny has slowed progress (European Commission, 2024). AI-driven systems in **quantitative trading** execute high-speed, precise trades, and generative AI is reshaping capital markets (EY, 2024). Additionally, **predictive analytics** forecast market trends for better decision-making (MIT Sloan, 2024), reinforcing earlier findings on AI's impact on data management and productivity (Poola, 2017; Brynjolfsson and McAfee, 2014).

C. Manufacturing Sector

AI is transforming manufacturing by optimizing processes, reducing costs, and improving quality. **Predictive maintenance** uses AI to analyze sensor data, anticipating machine failures and minimizing downtime (Appinventiv, 2025). **Computer vision systems** enhance quality control by detecting defects in real-time (IBM, 2024). AI-driven robots and **cobots** automate repetitive tasks, boosting efficiency and safety (Forbes, 2023). **Generative AI** accelerates product development through scenario modeling (AI Multiple, 2024). Furthermore, AI optimizes the supply chain by predicting demand and managing inventory, enabling real-time decision-making (TechTarget, 2024). These applications demonstrate AI's significant productivity boost in modern manufacturing (Brynjolfsson and McAfee, 2014).

D. Other Sectors

AI extends to legal research, predicting case outcomes (Andersen, 2023), and retail, personalizing customer experiences (Gurkaynak et al., 2016). It supports military planning (Poola, 2017; Wyatt Hoffman & Kim, 2023) and weather forecasting (Jatin Borana, 2016). Broader

applications are seen in cybersecurity and smart cities, with challenges in implementation (Sarker, 2021; Jiamin Yin et al., 2021)

Future Scope of AI

The future of AI is poised for significant growth, with several key trends shaping its trajectory. The **widespread adoption across industries** will see AI integrated into sectors like healthcare, finance, and manufacturing to enhance efficiency and innovation. According to Built In (2025), a substantial portion of enterprise businesses have already adopted or are considering AI, demonstrating its growing influence.

A. Advancements in Generative AI

Generative AI will become even more sophisticated, creating new content from text and images to code, revolutionizing creative and industrial processes. McKinsey (2024) highlights its role in industry-specific applications, suggesting that this technology is moving beyond simple content creation to become a core tool in various fields.

B. Efficient and Smaller AI Models

"**Edge AI**" and the development of smaller, more efficient models will make AI capabilities accessible on devices with limited resources. IBM (2024) notes the development of models like mini GPT 4o-mini, which will make AI more pervasive and cost-effective, enabling its use in everyday devices and applications.

C. Job Transformation

The rise of AI will inevitably lead to **job transformation**. While some jobs may be displaced, new roles will emerge that require creativity, empathy, and complex problem-solving. TechTarget (2024) predicted that human lawyers role may get reduced which underscores the need for workforce **reskilling**.

D. Ethical Development and Regulation

A deep focus on transparency, fairness, and accountability in AI systems is necessitated. The USC Annenberg (2024) discusses the ethical dilemmas and the evolving need for governance to ensure AI's responsible use.

E. Pursuit of Artificial General Intelligence (AGI)

The pursuit for **Artificial General Intelligence (AGI)** that can outperform any intellectual task a human can, will continue to drive research. Its not certain when it will be realized but the race to develop AGI highlights its potential to reshape society and raise significant questions about its implications (Pew, 2018).

Conclusion: Summary and Implications

This review highlights impact of AI across sectors and its ethical challenges, emphasizing the urgent need for strategic planning. AI paradoxically boosts efficiency but brings human autonomy at risk too, especially when corporate control can turn algorithms into manipulative tools. Reclaiming rationality demands **transparency**, **equitable governance**, and **ethical innovation**. As noted by O'Neil in *Weapons of Math Destruction* (2016), ungoverned AI is a threat to democracy. The conclusion is a synthesis of findings, acknowledging AI's ability to surpass bounded

rationality while introducing emerging issues like **bias** and **privacy**. It advocates for **hybrid systems** and **policy reforms**, grounded in recent studies and inspiring future engagement.

Summary and Implications

This expanded review underscores AI's transformative impact across sectors, its ethical challenges requiring urgent attention, and its future scope demanding strategic planning. The integration of recent studies ensures the paper reflects the state of AI as of March 24, 2025, providing a foundation for further research and policy development.

Table 1: Summary of Key Findings

Section	Key Findings	Example Citations
Bounded Rationality	Humans rely on satisficing and heuristics due to cognitive constraints like limited memory and processing speed.	Simon (1955), Kahneman & Tversky (1979), Gigerenzer & Goldstein (1996)
AI and Decision Making	AI supports decision-making across sectors but risks automation bias, overdependence, and reduced critical thinking.	Chen et al. (2021), Silver et al. (2016), Ruginski et al. (2019)
AI-Influenced Rationality	AI nudges behavior through design choices and can manipulate autonomy via dark patterns and addictive architectures.	Harris (2018), Zuboff (2019), Yeung (2017)
AI Applications	AI has revolutionized healthcare, finance, manufacturing, and education through diagnostics, forecasting, and automation.	Jones et al. (2023), Built In (2023), Rahmaniar et al. (2023)
Ethical Challenges	Bias in data, lack of transparency, and surveillance capitalism demand regulation and explainability in AI systems.	Buolamwini & Gebru (2018), UNESCO (2021), Arrieta et al. (2020)
Future Scope	Growth areas include explainable AI (XAI), quantum AI, edge computing, and AI literacy to mitigate ethical and access gaps.	OECD (2023), McKinsey (2024), Kairouz et al. (2021)
Hybrid AI-Human Systems	Collaborative frameworks can combine machine efficiency with human judgment for more resilient outcomes.	Chen et al. (2021), D'Acunto et al. (2019)
Environmental Impact	Training large AI models contributes to climate costs, disproportionately affecting low-income communities.	Strubell et al. (2019), Hogan (2021)
Corporate Control	AI under corporate dominance can exacerbate inequality and commodify behavior for profit-driven manipulation.	Rosenblat (2018), Khan (2017), Epstein & Robertson (2015)
Governance Solutions	Decentralized learning, algorithmic audits, and regulatory sandboxes offer equitable paths forward.	Cath et al. (2018), EU AI Act (2021), Kairouz et al. (2021)



This review illuminates that artificial intelligence and human decision making bear a complex relationship, a relationship that is both mutually beneficial yet to be ethically vigilant. The paper's synthesis of more than 100 sources reveals pressing research gaps. Longitudinal studies are needed to assess AI's long-term cognitive effects: does GPS reliance atrophy spatial memory (Ruginski et al., 2019)? Can participatory design frameworks like "algorithmic citizen juries" mitigate bias in predictive policing?

For policymakers, three imperatives stand out:

1. Expanding AI literacy to combat automation bias.
2. Enforcing algorithmic accountability through third-party audits.
3. Redistributing AI's economic gains via data dividends and reskilling programs.

Ultimately, this review affirms that intelligence—artificial or human—is defined not by its bounds but by its ethical ends. As Simon (1990) cautioned, "The capacity of the human mind is very small compared to the size of the problems whose solution is required for objectively rational behavior." In navigating AI's future, humility, equity, and interdisciplinary rigor must guide our path.

References

- Alter, A. (2017). *Irresistible: The rise of addictive technology and the business of keeping us hooked*. Penguin.
- Andersen. (2023). *Beyond the hype: Real-world applications of AI*. [Executive summary].
- Anderson, J. R. (1990). *The adaptive character of thought*. Psychology Press.
- Anderson, J. R., & Schooler, L. J. (1991). Reflections of the environment in memory. *Psychological Science*, 2(6), 396–408.
- Angwin, J., Larson, J., Mattu, S., & Kirchner, L. (2016). *Machine bias*. ProPublica.
- Arrieta, A. B., et al. (2020). Explainable artificial intelligence (XAI): Concepts, taxonomies, opportunities, and challenges toward responsible AI. *Information Fusion*, 58, 82–115.
- Awad, E., et al. (2018). The Moral Machine experiment. *Nature*, 563(7729), 59–64.
- Baeza-Yates, R. (2022). Ethical challenges in AI. In *Proceedings of the Fifteenth ACM International Conference on Web Search and Data Mining* (pp. 1–2). Association for Computing Machinery. <https://doi.org/10.1145/3488560.3498370>
- Baker, R. S., & Inventado, P. S. (2014). Educational data mining and learning analytics. In J. A. Larusson & B. White (Eds.), *Learning analytics: From research to practice*. Springer.
- Balakrishnan, A., et al. (2020). AI-driven traffic management in smart cities. *IEEE Transactions on Intelligent Transportation Systems*, 21(8), 3154–3165.
- Beck, J., Stern, M., & Haugsjaa, E. (1996). Applications of AI in education. *XRDS: Crossroads, The ACM Magazine for Students*, 3(1), 11–15.
- Bender, E. M., et al. (2021). On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 610–623).
- Benjamin, R. (2019). *Race after technology: Abolitionist tools for the New Jim Code*. Polity Press.
- Biamonte, J., et al. (2017). Quantum machine learning. *Nature*, 549(7671), 195–202.
- Borana, J. (2016). Applications of artificial intelligence and associated technologies. In *Proceedings of the International Conference on Emerging Technologies in Engineering, Biomedical, Management and Science (ETEBMS-2016)* (pp. 64–68).
- Borenstein, J., & Howard, A. (2021). Emerging challenges in AI and the need for AI ethics education. *AI and Ethics*, 1, 61–65. <https://doi.org/10.1007/s43681-020-00002-7>
- Breiman, L., Friedman, J. H., Olshen, R. A., & Stone, C. J. (1984). *Classification and regression trees*. Routledge.
- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32.
- Bridle, J. (2020). *New dark age: Technology and the end of the future*. Verso.
- Brown, T. B., et al. (2020). Language models are few-shot learners. In *Advances in Neural Information Processing Systems* (Vol. 33, pp. 1877–1901).
- Brynjolfsson, E., & McAfee, A. (2014). *The second machine age: Work, progress, and prosperity in a time of brilliant technologies*. W. W. Norton & Company.
- Buolamwini, J., & Gebru, T. (2018). Gender shades: Intersectional accuracy disparities in commercial gender classification. In *Proceedings of Machine Learning Research* (Vol. 81, pp. 1–15).
- Campbell, M., Hoane, A. J., & Hsu, F.-h. (2002). Deep Blue. *Artificial Intelligence*, 134(1–2), 57–83.
- Campolo, A., et al. (2017). *AI Now 2017 report*. AI Now Institute.
- Carr, N. (2014). *The glass cage: Automation and us*. W. W. Norton & Company.
- Cath, C., et al. (2018). Governing artificial intelligence: Ethical, legal, and technical opportunities and challenges. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 376(2133), 20180080.
- Chen, J. H., et al. (2021). Human–AI collaboration in clinical decision-making. *The New England Journal of Medicine*, 385(14), 1345–1354.
- Chen, L., & Sheldon, M. (2016). Dynamic pricing in a labor market: Surge pricing and flexible work on the Uber platform. In *Proceedings of the 2016 ACM Conference on Economics and Computation* (pp. 455–456).
- Chesney, R., & Citron, D. K. (2019). Deep fakes: A looming challenge for privacy, democracy, and national security. *California Law Review*, 107(6), 1753–1820.
- Cortes, C., & Vapnik, V. (1995). Support-vector networks. *Machine Learning*, 20(3), 273–297.
- Crevier, D. (1993). *AI: The tumultuous history of the search for artificial intelligence*. Basic Books.
- Cummings, M. L. (2004). Automation bias in intelligent time-critical decision support systems. In *AIAA 1st Intelligent Systems Technical Conference*. American Institute of Aeronautics and Astronautics.
- Cyert, R. M., & March, J. G. (1963). *A behavioral theory of the firm*. Prentice-Hall.
- D'Acunto, F., et al. (2019). Robo-advising. *Journal of Financial Economics*, 134(2), 443–465.
- Dal Pozzolo, A., et al. (2015). Calibrating probability with undersampling for unbalanced classification. In *2015 IEEE Symposium Series on Computational Intelligence*. IEEE.
- Damasio, A. R. (1994). *Descartes' error: Emotion, reason, and the human brain*. Putnam.
- Dastin, J. (2018). Amazon scraps secret AI recruiting tool that showed bias against women. *Reuters*.
- Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard Business Review*, 96(1), 108–116.
- Davies, M., et al. (2021). Loihi 2: A neuromorphic chip for edge AI. *IEEE Micro*, 41(5), 82–99.
- De Fauw, J., et al. (2018). Clinically applicable deep learning for diagnosis and referral in retinal disease. *Nature Medicine*, 24(9), 1342–1350.
- Dwork, C., et al. (2006). Calibrating noise to sensitivity in private data analysis. In *Theory of Cryptography Conference*.

- Easley, D., et al. (2021). High-frequency trading and market stability. *Journal of Financial Markets*, 54, 100596.
- Epstein, R., & Robertson, R. E. (2015). The search engine manipulation effect (SEME) and its possible impact on the outcomes of elections. *Proceedings of the National Academy of Sciences of the United States of America*, 112(33), E4512–E4521.
- European Commission. (2021). *Proposal for a regulation laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)*. EUR-Lex.
- Evans, R., & Gao, J. (2016). DeepMind AI reduces energy used for cooling Google data centers by 40%. *Google Blog*.
- Feigenbaum, E. A. (1977). The art of artificial intelligence: Themes and case studies of knowledge engineering. In *Proceedings of the Fifth International Joint Conference on Artificial Intelligence* (pp. 1014–1029).
- Floridi, L. (2014). *The 4th revolution: How the infosphere is reshaping human reality*. Oxford University Press.
- Gebru, T., et al. (2021). Datasheets for datasets. *Communications of the ACM*, 64(12), 86–92.
- Gigerenzer, G. (2008). Why heuristics work. *Perspectives on Psychological Science*, 3(1), 20–29.
- Gigerenzer, G., & Gaissmaier, W. (2011). Heuristic decision making. *Annual Review of Psychology*, 62, 451–482.
- Gigerenzer, G., & Goldstein, D. G. (1996). Reasoning the fast and frugal way: Models of bounded rationality. *Psychological Review*, 103(4), 650–669.
- Goldstein, D. G., & Gigerenzer, G. (2002). Models of ecological rationality: The recognition heuristic. *Psychological Review*, 109(1), 75–90.
- Gomez-Urbe, C. A., & Hunt, N. (2015). The Netflix recommender system: Algorithms, business value, and innovation. *ACM Transactions on Management Information Systems*, 6(4), Article 13, 1–19.
- Goodfellow, I. J., et al. (2015). Explaining and harnessing adversarial examples. In *International Conference on Learning Representations*.
- Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep learning*. MIT Press.
- Griffiths, T. L., et al. (2015). Rational use of cognitive resources: Levels of analysis between the computational and the algorithmic. *Topics in Cognitive Science*, 7(2), 217–229.
- Gupta, A., Dwivedi, D. N., & Shah, J. (2023). Ethical challenges for AI-based applications. In *Artificial intelligence applications in banking and financial services*. Springer. https://doi.org/10.1007/978-981-99-2571-1_10
- Haig, M. (2021). *The Facebook files: A Wall Street Journal investigation*.
- Haleem, A., Javaid, M., & Khan, I. H. (2019). Current status and applications of artificial intelligence (AI) in the medical field: An overview. *Current Medicine Research and Practice*, 9(6), 231–237. <https://doi.org/10.1016/j.cmrp.2019.11.005>
- Hare, T. A., et al. (2011). Interactions between affect and decision-making during goal pursuit. *Nature Neuroscience*, 14(6), 720–726.
- Harris, T. (2018). *How technology hijacks people's minds*.
- Hastie, T., Tibshirani, R., & Friedman, J. (2009). *The elements of statistical learning: Data mining, inference, and prediction* (2nd ed.). Springer.
- Henrich, J., et al. (2001). In search of Homo economicus: Behavioral experiments in 15 small-scale societies. *American Economic Review*, 91(2), 73–78.
- Hill, K. (2021). The secretive company that might end privacy as we know it. *The New York Times*.
- Hoffman, W., & Kim, H. M. (2023). *Reducing the risks of artificial intelligence for military decision advantage*. Center for Security and Emerging Technology. <https://doi.org/10.51593/2021CA008>
- Hogan, M. (2021). Data flows and water woes: The environmental costs of cloud computing. *Environmental Humanities*, 13(1), 168–185.
- Holzinger, A., et al. (2019). Interactive machine learning for health informatics. *Artificial Intelligence in Medicine*, 94, 1–5.
- Hosanagar, K., et al. (2014). Will the global village fracture into tribes? Recommender systems and their effects on consumer fragmentation. *Management Science*, 60(4), 805–823.
- Huang, L., et al. (2021). Adversarial attacks on financial time series. In *Proceedings of the ACM SIGKDD Conference on Knowledge Discovery and Data Mining*.
- Hussain, F., & Qamar, U. (2016). Identification and correction of misspelled drug names in electronic medical records (EMR). In *Proceedings of the 18th International Conference on Enterprise Information Systems* (Vol. 2, pp. 333–338).
- Hutchins, W. J. (1986). *Machine translation: Past, present, future*. Ellis Horwood.
- Institute of Electrical and Electronics Engineers. (2019). *Ethically aligned design: A vision for prioritizing human well-being with autonomous and intelligent systems* (1st ed.). IEEE Standards Association.
- Jacowitz, K. E., & Kahneman, D. (1995). Measures of anchoring in estimation tasks. *Personality and Social Psychology Bulletin*, 21(11), 1161–1166.
- Khan, L. M. (2017). Amazon's antitrust paradox. *Yale Law Journal*, 126(3), 710–805.
- Kahneman, D. (2011). *Thinking, fast and slow*. Farrar, Straus and Giroux.
- Kahneman, D., Knetsch, J. L., & Thaler, R. H. (1990). Experimental tests of the endowment effect and the Coase theorem. *Journal of Political Economy*, 98(6), 1325–1348.
- Kahneman, D., Slovic, P., & Tversky, A. (Eds.). (1982). *Judgment under uncertainty: Heuristics and biases*. Cambridge University Press.
- Kahneman, D., & Tversky, A. (1979). Prospect theory: An analysis of decision under risk. *Econometrica*, 47(2), 263–292.
- Kairouz, P., et al. (2021). Advances and open problems in federated learning. *Foundations and Trends in Machine Learning*, 14(1–2), 1–210.
- Kalusivalingam, A. K. (2018). Ethical considerations in AI: Historical perspectives and contemporary challenges. *Journal of Innovative Technologies*, 1(1), 1–8.
- Khan, S., Farid, S., Bhat, M. S., Hamid, U., Dar, B. A., & Dar, M. A. (2022). The role and importance of artificial intelligence in today's scenario. *Journal of the Maharaja Sayajirao University of Baroda*, 56(1[8]).
- Kirilenko, A., et al. (2017). The flash crash: High-frequency trading in an electronic market. *The Journal of Finance*, 72(3), 967–998.
- Krizhevsky, A., Sutskever, I., & Hinton, G. E. (2012). ImageNet classification with deep convolutional neural networks. In *Advances in Neural Information Processing Systems* (Vol. 25, pp. 1097–1105).
- Lakhani, P., et al. (2020). Machine learning in radiology: Applications beyond image interpretation. *Journal of the American College of Radiology*, 17(7), 977–984.
- Lebedev, M. A., & Nicolelis, M. A. L. (2017). Brain-machine interfaces: From basic science to neuroprostheses and neurorehabilitation. *Physiological Reviews*, 97(2), 767–837.
- Lee, M. K., et al. (2015). Working with machines: The impact of algorithmic and data-driven management on human workers. In *Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems* (pp. 1603–1612).
- Levin, K., et al. (2012). Overcoming the tragedy of super wicked problems: Constraining our future selves to ameliorate global climate change. *Policy Sciences*, 45(2), 123–152.
- Lewis, R. L., et al. (2014). Computational rationality: A converging paradigm for intelligence in brains, minds, and machines. *Science*, 349(6245), 273–278.
- Li, F., Ruijs, N., & Lu, Y. (2023). Ethics and AI: A systematic review on ethical concerns and related strategies for designing with AI in healthcare. *AI*, 4(1), 28–53. <https://doi.org/10.3390/ai4010003>



- Lighthill, J. (1973). Artificial intelligence: A general survey. In *Artificial intelligence: A paper symposium*. Science Research Council.
- Lindsay, R. K., et al. (1980). *Applications of artificial intelligence for organic chemistry: The DENDRAL project*. McGraw-Hill.
- Luguri, J., & Strahilevitz, L. J. (2021). Shining a light on dark patterns. *Journal of Legal Analysis*, 13(1), 43–109.
- Lum, K., & Isaac, W. (2016). To predict and serve? *Significance*, 13(5), 14–19.
- Machado, H., Silva, S., & Neiva, L. (2025). Publics' views on ethical challenges of artificial intelligence: A scoping review. *AI and Ethics*, 5, 139–167. <https://doi.org/10.1007/s43681-023-00387-1>
- March, J. G. (1991). Exploration and exploitation in organizational learning. *Organization Science*, 2(1), 71–87.
- Marr, B. (2018). How is AI used in education? Real-world examples of today and a peek into the future. *Forbes*.
- Mathur, A., et al. (2021). Dark patterns at scale: Findings from a crawl of 11K shopping websites. *Proceedings of the ACM on Human-Computer Interaction*, 3(CSCW), 1–32.
- Matthias, A. (2004). The responsibility gap: Ascribing responsibility for the actions of learning automata. *Ethics and Information Technology*, 6(3), 175–183.
- McKay, F., Williams, B. J., Prestwich, G., Bansal, D., Hallowell, N., & Treanor, D. (2022). The ethical challenges of artificial intelligence-driven digital pathology. *The Journal of Pathology: Clinical Research*. Advance online publication. <https://doi.org/10.1002/cjp.2.263>
- Mehrabi, N., et al. (2021). A survey on bias and fairness in machine learning. *ACM Computing Surveys*, 54(6), Article 115, 1–35.
- Miller, G. A. (1956). The magical number seven, plus or minus two: Some limits on our capacity for processing information. *Psychological Review*, 63(2), 81–97.
- Mnih, V., et al. (2015). Human-level control through deep reinforcement learning. *Nature*, 518(7540), 529–533.
- Newell, A., & Simon, H. A. (1956). The logic theory machine. *IRE Transactions on Information Theory*, 2(3), 61–79.
- Ngiam, K. Y., & Teo, H. H. (2021). The role of artificial intelligence applications in real-life clinical practice: Systematic review. *Journal of Medical Internet Research*, 23(4), Article e21359.
- National Highway Traffic Safety Administration. (2023). *Investigation into Tesla Autopilot crashes*.
- Nyholm, S. (2018). The ethics of crashes with self-driving cars: A roadmap. *Philosophy Compass*, 13(7), e12507.
- O'Neil, C. (2016). *Weapons of math destruction: How big data increases inequality and threatens democracy*. Crown.
- Obermeyer, Z., et al. (2019). Dissecting racial bias in an algorithm used to manage the health of populations. *Science*, 366(6464), 447–453.
- OpenAI. (2023). *GPT-4 technical report*. arXiv. <https://doi.org/10.48550/arXiv.2303.08774>
- Orús, R., et al. (2019). Quantum computing for finance: Overview and prospects. *Reviews in Physics*, 4, 100028.
- Parasuraman, R., & Manzey, D. H. (2010). Complacency and bias in human use of automation: An attentional integration. *Human Factors*, 52(3), 381–410.
- Pariser, E. (2011). *The filter bubble: What the Internet is hiding from you*. Penguin Press.
- Patterson, D., et al. (2021). Carbon emissions and large neural network training. arXiv. <https://doi.org/10.48550/arXiv.2104.10350>
- Payne, J. W., et al. (1993). *The adaptive decision maker*. Cambridge University Press.
- Pedro, F., Subosa, M., Rivas, A., & Valverde, P. (2019). *Artificial intelligence in education: Challenges and opportunities for sustainable development*. UNESCO.
- Pentland, A. (2014). *Social physics: How good ideas spread—The lessons from a new science*. Penguin Press.
- Poola, I. (2017). How artificial intelligence is impacting real life every day. *International Journal for Advance Research and Development*, 2(10), 96–100.
- Preskill, J. (2018). Quantum computing in the NISQ era and beyond. *Quantum*, 2, Article 79.
- Rahmaniar, W., Maarif, A., ul Haq, Q. M., & Iskandar, M. E. (2023). *AI in industry: Real-world applications and case studies*. Authorea Preprints.
- Rahmaniar, W., Maarif, A., ul Haq, Q. M., et al. (2023). *AI in industry: Real-world applications and case studies*. TechRxiv. <https://doi.org/10.36227/techrxiv.23993565.v1>
- Rahwan, I., et al. (2019). Machine behaviour. *Nature*, 568(7753), 477–486.
- Raji, I. D., et al. (2022). AI literacy: A critical framework to understand the power of AI. In *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency* (pp. 1–10).
- Rajpurkar, P., et al. (2022). AI in radiology: Review of current literature. *Radiology*, 302(1), 1–12.
- Ribeiro, M. T., et al. (2016). 'Why should I trust you?' Explaining the predictions of any classifier. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 1135–1144).
- Rombach, R., et al. (2022). High-resolution image synthesis with latent diffusion models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (pp. 10684–10695).
- Rosenblat, A. (2018). *Uberland: How algorithms are rewriting the rules of work*. University of California Press.
- Rosenblat, A., & Stark, L. (2016). Algorithmic labor and information asymmetries: A case study of Uber's drivers. *International Journal of Communication*, 10, 3758–3784.
- Ross, C., & Swetlitz, I. (2018). IBM's Watson supercomputer recommended 'unsafe and incorrect' cancer treatments, internal documents show. *STAT*.
- Sanchez, T. W., Brenman, M., & Ye, X. (2024). The ethical concerns of artificial intelligence in urban planning. *Journal of the American Planning Association*, 91(2), 294–307. <https://doi.org/10.1080/01944363.2024.2355305>
- Sarker, I. H. (2021). Machine learning: Algorithms, real-world applications and research directions. *SN Computer Science*, 2, Article 160. <https://doi.org/10.1007/s42979-021-00592-x>
- Slimi, Z., & Carballido, B. V. (2023). Navigating the ethical challenges of artificial intelligence in higher education: An analysis of seven global AI ethics policies. *TEM Journal*, 12(2), 590–602.
- Slimi, Z., & Carballido, B. V. (2023). Systematic review: AI's impact on higher education—Learning, teaching, and career opportunities. *TEM Journal*, 12(3), 1627–1637.
- Stahl, B. C. (2021). Ethical issues of AI. In *Artificial intelligence for a better future: An ecosystem perspective on the ethics of AI and emerging digital technologies*. Springer. https://doi.org/10.1007/978-3-030-69978-9_4
- Sun, W., Li, C., & Ren, R. (2019). Artificial intelligence and its application in computer network technology. *Journal of Physics: Conference Series*, 1237, Article 022142. <https://doi.org/10.1088/1742-6596/1237/2/022142>
- Usmani, U. A., Happonen, A., & Watada, J. (2022). Enhancing artificial intelligence control mechanisms: Current practices, real-life applications, and future views. In *Proceedings of the Future Technologies Conference* (pp. 287–306). Springer International Publishing.
- Zhai, X., Chu, X., Chai, C. S., Jong, M. S.-Y., Istenic, A., Spector, M., et al. (2021). A review of artificial intelligence (AI) in education from 2010 to 2020. *Complexity*, 2021, Article 8812542.
- Zhang, Y., Balochian, S., Agarwal, P., Bhatnagar, V., & Housheya, O. J. (2014). Artificial intelligence and its applications. *Mathematical Problems in Engineering*, 2014

GJEIS Prevent Plagiarism in Publication

DELNET-Developing Library Network, New Delhi in collaboration with BIPL has launched “DrillBit : Plagiarism Detection Software for Academic Integrity” for the member institutions of DELNET. It is a sophisticated plagiarism detection software which is currently used by 700+ Institutions in India and outside. DrillBit is a global checker that uses the most advanced technology to catch the most sophisticated forms of plagiarism, plays a critical function for students and instructors and tag on a fully-automatic machine learning text- recognition system made for detecting, preventing and handling plagiarism and trusted by thousands of institutions across worldwide. DrillBit - Plagiarism Detection Software has been preferred for empanelment with AICTE and NEAT 3.0 (National Education Alliance for Technology) and contributing towards enhanced learning outcomes in India. On the other hand software uses a number of methods to detect AI-generated content, including, checking for repetitive phrases or sentences and AI-generated writing. As part of a larger global organization GJEIS (www.gjeis.com) and DrillBit better equipped to anticipate the foster an environment of academic integrity for educators and students around the globe. DrillBit is GDPR compliant with privacy by design and an uptime of 99.9% and have trust to be the partner in academic integrity (<https://www.drillbitplagiarism.com>) tool to check the originality and further affixed the similarity index which is {01%} in this case (See below Annexure 18.1.8). Thus, the reviewers and editors are of view to find it suitable to publish in this Volume 18, Issue-1, Jan-Mar 2026.

Annexure 18.1.8

Submission Date	Submission Id	Word Count	Character Count
27-Jan-2026	6119111 (DrillBit)	7913	56972

Analyzed Document	Submitter email	Submitted by	Similarity
5.3 RoL3_ Krishnendu_GJEIS Jan-Mar 2026.docx	krishnendudas20@gmail.com	Krishnendu Das	01%



DrillBit Similarity Report

1	4	A
SIMILARITY %	MATCHED SOURCES	GRADE
LOCATION	MATCHED DOMAIN	%
1	www.humanrightsresearch.org	<1

A.Satisfactory (0-10%)
B.Upgrade (11-40%)
C.Poor (41-60%)
D.Unacceptable (61-100%)

3	rsisinternational.org	<1	Publication
11	arxiv.org	<1	Publication
13	arxiv.org	<1	Publication

Reviewers Memorandum

Reviewer's Comment 1: The manuscript is well written and presents the research topic in a clear, logical, and coherent manner. The overall organization of the paper enables readers to follow the arguments with ease, and the writing quality is satisfactory. The study addresses a relevant research problem, and the objectives are generally aligned with the scope of the research. To further enhance the quality of the manuscript, the research gap should be articulated more explicitly to provide a stronger justification for the study. Additionally, the research objectives could be stated more clearly to improve their precision and ensure a better understanding of the intended outcomes of the research.

Reviewer's Comment 2: The manuscript demonstrates a good understanding of the research topic, and the concepts have been defined clearly and systematically, making the paper easy to comprehend. The introduction provides a suitable background and establishes the context of the study. However, the research gaps should be more explicitly linked to the discussion presented in the preceding literature review to strengthen the logical progression of the manuscript. Furthermore, greater emphasis on recent and relevant literature in the introduction would improve the theoretical foundation of the study and better demonstrate its significance within the current body of knowledge.

Reviewer's Comment 3: The manuscript presents a relevant and timely research topic and is written in a clear, organized, and professional manner. The overall structure of the paper is coherent, and the discussion reflects a sound understanding of the subject. The study has the potential to make a meaningful contribution to the literature. Nevertheless, the literature review could be expanded to include a more comprehensive and critical discussion of previous studies. In addition, the methodology section would benefit from a more detailed description of the research procedures and analytical approach, which would improve the transparency, rigor, and reproducibility of the research.



Krishnendu Das and Anamika Choudhary
“Bounded Rationality and AI in Decision Making”
Volume-18, Issue 1, Jan-Mar 2026. (www.gjeis.com)

<https://doi.org/10.18311/gjeis/2026>
Volume-18, Issue 1, Jan-Mar 2026

Online iSSN : 0975-1432, Print iSSN : 0975-153X
Frequency : Quarterly, Published Since : 2009

Google Citations: Since 2009
H-Index = 96
i10-Index: 964

Source: <https://scholar.google.co.in/citations?user=S47TtNkAAAAJ&hl=en>

Conflict of Interest: Author of a Paper
had no conflict neither financially nor academically.

**Editorial
Excerpt**

The article has 1% of plagiarism which is the accepted percentage as per the norms and standards of the journal for publication. As per the editorial board's observations and blind reviewers' remarks the paper had some minor revisions which were communicated on a timely basis to the authors (Krishnendu and Anamika), and accordingly, all the corrections had been incorporated as and when directed and required to do so. The comments related to this manuscript are noticeably related to the theme "BOUNDED RATIONALITY AND AI IN DECISION MAKING" both subject-wise and research-wise. The manuscript examines the complex interplay between human decision-making and artificial intelligence through the perspective of Herbert Simon's theory of bounded rationality. It explores how AI can both enhance and constrain human cognitive processes while critically addressing the ethical, environmental, and governance implications associated with its increasing adoption. The study investigates whether AI truly overcomes the inherent limitations of human decision-making or merely replaces existing cognitive constraints with new challenges related to algorithmic bias, accountability, sustainability, and governance. In doing so, the paper aims to provide a balanced understanding of AI's role in shaping contemporary decision-making processes. After comprehensive reviews and the editorial board's remarks, the manuscript has been categorized and decided to publish under the "**Literature Review Paper**" category.

Acknowledgement

The acknowledgment section is an essential part of all academic research papers. It provides appropriate recognition to all contributors for their hard work and effort taken while writing a paper. The data presented and analyzed in this paper by (Krishnendu and Anamika) were collected first handily and wherever it has been taken the proper acknowledgment and endorsement depicts. The authors are highly indebted to others who facilitated accomplishing the research. Last but not least, endorse all reviewers and editors of GJEIS in publishing in the present issue.

Disclaimer

All views expressed in this paper are my/our own. Some of the content is taken from open-source websites & some are copyright free for the purpose of disseminating knowledge. Those some we/I had mentioned above in the references section and acknowledged/cited as when and where required. The author/s have cited their joint own work mostly, and tables/data from other referenced sources in this particular paper with the narrative & endorsement have been presented within quotes and reference at the bottom of the article. Accordingly & appropriately. Finally, some of the contents are taken or overlapped from open-source Websites for knowledge purposes. Those some of i/we had mentioned above in the references section. On the other hand, opinions expressed in this paper are those of the author and do not reflect the views of the GJEIS. The authors have made every effort to ensure that the information in this paper is correct, any remaining errors and deficiencies are solely their responsibility.