

AI in Cloud Platforms: Ethical Considerations

Ali M. Iqbal¹; Majed Al Otaibi²; Khalid Aljaghthami³

^{1,2,3}Enterprise Digital Solutions Division, Saudi Aramco, Dhahran, Saudi Arabia

Publication Date: 2025/11/08

Abstract: Cloud computing AI is revolutionizing software development, security, and operations. Among those breakthroughs is AI-generated code — machine-authored logic that automates processes during the lifecycle of the cloud. In addition to being a boon to the growth of lean and resource-efficient systems, automation can open up more complex ethical dilemmas. These encompass decision logic bias, threat detection opacity, a lack of accountability in automated remediation, and the privacy risks posed by data-driven personalization. In this work we consider the ethical issues surrounding AI code generation for cloud platforms in three domains: cloud development, security operations, and decision-making systems. By applying some specific technical examples and combining a synthesis of recent literature we argue the extent to which AI has the power and potential for threats to be realized in cloud-based environments. We are calling for governance mechanisms such as fairness-aware models, explainable AI, human-in-the-loop supervision, and consent-aware data practices. The findings of our study indicate that ethical AI is not an option, but rather a requirement, for a secure, transparent and accountable cloud platform that should be provided.

Keywords: Artificial Intelligence, Cloud Computing, Ethics, Automation, Security, Decision-Making, Governance.

How to Cite: Ali M. Iqbal; Majed Al Otaibi; Khalid Aljaghthami (2025) AI in Cloud Platforms: Ethical Considerations. *International Journal of Innovative Science and Research Technology*, 10(11), 85-90 <https://doi.org/10.38124/ijisrt/25nov133>

I. INTRODUCTION

Integration of Artificial Intelligence (AI) in cloud computing has changed the way software is created, deployed, and maintained. Among the most disruptive developments is the emergence of AI-generated code – machine-generated computer scripts, configurations, and logic that automate work that was previously done by humans. AI-generated code is embedded across the cloud software lifecycle (from infrastructure provisioning to security orchestration and continuous deployment).

The benefits of this automation are tremendous from speed and scalability to cost and efficiency but it also presents challenging ethical dilemmas. AI systems may produce code based on biases; it also lacks transparency or could create security issues. And, when code is written by non-human writers, the issues of accountability and intellectual property become increasingly blurred.

In this paper we explore the ethical challenges of AI-generated code in cloud, and particularly on three important domains: (1) AI-driven automation in cloud development; (2) AI-booster security in multi-cloud environments; and (3) AI-guided decision-making in cloud operations. Taking technical cases and incorporating a review of recent work, we discuss how such systems can both contribute to and threaten modern clouds. The aim is to illustrate the necessary requirement for responsible AI practices, governance models, and human

supervision if automation's benefits are to be achieved without ethical cost.

II. LITERATURE REVIEW

Recent literature identifies ethical issues related to AI-derived code in cloud platforms. To prevent discriminatory outcomes in cloud-based analytics, Vayyala proposes fairness-aware algorithms and governance frameworks [1]. Floridi and Cowls propose accountability models that assign responsibility for AI decisions, highlighting the importance of human oversight of automated systems. Surya et al. show that CNNs, LSTMs and similar deep learning models have been successful in threat detection, achieving over 93% accuracy and precision [2]. Singh and Choudhry recommend self-healing and predictive maintenance procedures, and warn against ethical gaps in AI-guided decision-making [3]. Overuse of AI in public clouds is a concern for Linthicum, warning that enterprises should deploy AI appropriately in the cloud where businesses must match ROI and operational relevance [4]. Polamarasetti et al. demonstrate that time-series models such as ARIMA optimize resource allocation and alleviate performance bottlenecks [5]. Including transparency, fairness and human oversight in cloud-based AI systems, these studies show.

➤ *AI-Powered Automation in Cloud Development*

- *Infrastructure Management*

Through workload patterns and demand forecasting, AI can optimize cloud resource allocation. For example, by examining historical traffic data, an AI model can predict peak usage periods and automatically provision the appropriate server capacity to maintain performance and cost-efficiency.

AI also improves cost efficiencies by spotting underutilized virtual machines and directing shutdowns for off-hours. One

other positive is predictive maintenance, in which AI forecasts hardware breakdowns based on data trends with an emphasis on preemptive maintenance, ensuring that downtime is minimized [4]

- *Code Generation*

AI models assist developers by generating code snippets, reducing manual effort and accelerating software development. These tools can analyze large datasets, generate boilerplate code, and optimize algorithms for specific problems. For example:

AI suggests a function to check if a number is prime
<pre>def is_prime(n): if n <= 1: return False for i in range(2, int(n**0.5) + 1): if n % i == 0: return False return True</pre>

However, ethical concerns arise when AI-generated code embeds insecure or biased logic. Consider this flawed password validation

Password validation
<pre>def validate_password(pw): return len(pw) > 8 or "password" in pw.lower()</pre>

Such patterns can introduce vulnerabilities if not reviewed by human developers [3].

While helpful, such suggestions may embed subtle biases or insecure logic if not reviewed [3]. Developers must validate AI-generated code to ensure ethical and secure outcomes.

- *Continuous Integration & Development (CI/CD)*

AI automates tasks of testing, debugging, and deploying software for higher reliability and less manual workload. It forecasts failure points, proposing test cases to check new code. If for example a new feature modifies a login module, AI can prioritize relevant security and authentication tests [5].

AI identifies optimal times to deploy updates, minimizing disruptions by scheduling them during periods of low traffic. . Singh and Choudhry emphasize that Gen AI can trigger self-healing behaviors, like restarting failed services or reallocating workloads, independently of humans [3].

➤ *AI-Enhanced Security in Cloud Platforms*

- *Threat and Anomaly Detection*

Machine learning methods keep track of logs and system metrics, also monitoring logs and system metrics to uncover anomalies, like an unexpected hike in login fails, which could mean that brute-force attacks take place [6].

AI-powered security systems analyze patterns to identify and mitigate cyber threats in real time. It can spot unusual patterns in network traffic, API usage, and security logs analysis in real-time [7]. Surya et al. demonstrate that Convolution Neural Networks (CNNs) and Long Term Memory Networks (LSTMs) outperform traditional methods in detecting anomalies and zero-day attacks, achieving over 93% accuracy and precision [2].

Algorithms for machine learning are deployed to monitor logs and system metrics to detect unusual activity. For instance, a spike in login failures might be a sign of brute-force attacks. These models dynamically evolve according to the changing threats [3].



Fig 1 AI-Enhanced Security in Cloud Platforms

- *Zero-Day Detection and Adaptive Defenses*

Machine learning models trained on behavioral signatures and threat intelligence can detect previously unknown vulnerabilities, known as zero-day exploits, before they are weaponized [6]. Hybrid training strategies using datasets like CSECIC-IDS2018 and KDD Cup 99 enable AI systems to recognize novel attack vectors and apply adaptive defenses before exploitation occurs [2].

- *Automated Incident Response*

Security Orchestration, Automation, and Response platforms use AI to trigger playbooks, contain threats, and initiate remediation across multi-cloud environments [8].

- *Identity Security and Zero Trust*

AI enables continuous identity verification across human and machine actors, enforcing zero-trust principles. It monitors behavioral patterns to detect anomalies in authentication, authorization, and privilege escalation [9].

- *AI Ops for Security Resilience*

AI for IT Operations (AIOps) correlates telemetry data across infrastructure layers to predict failures, detect faults, and maintain availability [10][11]. AI is trained on the data from telemetry across all the components of IT infrastructure that are used for predicting failures, detecting faults, and maintaining availability. Surya et al. demonstrate in a more practical context that AIOps-based deep learning frameworks can proactively detect system degradation and take remedial actions immediately [2].

- *LLM-Assisted Security Operations*

Large language models help Security Operations Centers to summarize alerts, prioritize incidents, and generate remediation steps [8][12]. They enhance analyst efficiency and minimize cognitive load. Gen AI interprets logs and metrics to identify root causes and suggest remediation steps, as noted by Singh and Choudhry [3].

- *Data Protection and Access Governance*

AI helps classify data faster, keeps track of access behaviors, and applies ongoing policy mechanisms to prevent unauthorized data exposure [7][8]. In multi-cloud settings, AI-based governance systems reduce data exfiltration risks by dynamically adjusting access policies based on usage context [8].

- *Automated Compliance Monitoring*

AI facilitates compliance with automated evidence collection, control validation, and reporting. Choudhry and Singh further stress that Gen AI can help maintain regulatory alignment across hybrid cloud ecosystems while adapting to evolving standards [3]. Also, automated compliance tools help organizations efficiently regulatory requirements [12].

- *AI-Driven Decision Making in Cloud Computing Platforms*

- *Predictive Analytics and Resource Allocation*

AI models, including ARIMA and LSTM, are used in cloud context to predict demand trends and prioritize resource utilization [3]. Such predictive models scale flexibly and use resources efficiently. Still, in the presence of biased training data, i.e. underrepresentation of specific user groups, these models may entrench unfair allocation and exclusion and even have service disparities [13]. Fairness and transparency with respect to predictive analytics is therefore essential for equitable management of cloud resources.

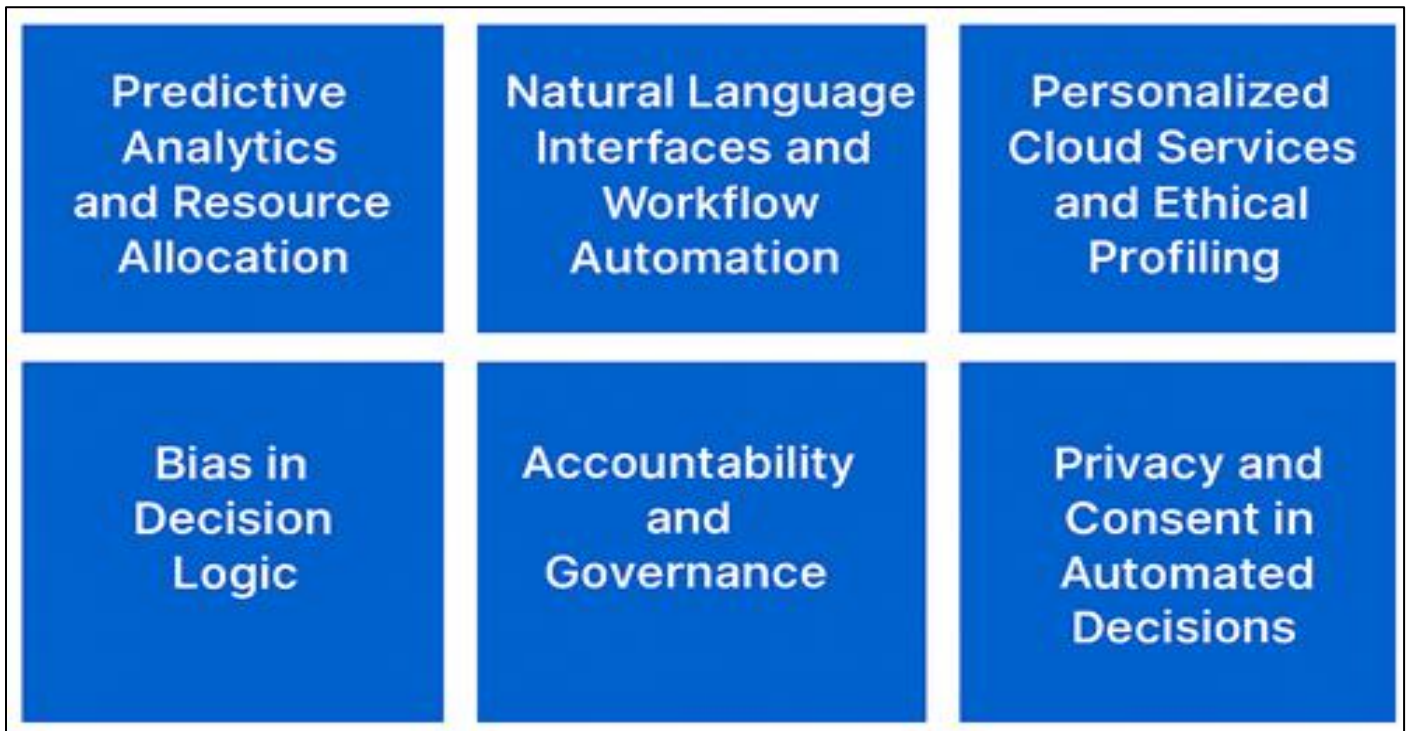


Fig 2 AI-Driven Decision Making in Cloud Platform

- *Natural Language Interfaces and Workflow Automation*

Large Language Models (LLMs) generate conversational interfaces that allow administrators and users to communicate with the cloud operating systems using human natural language commands. Singh and Choudhry assert that in order to keep safe from error and guarantee reliable automation, validation layers and verification protocols should be included in these conversational workflows [3]. Defend the above-mentioned safeguards from possible misinterpretation or against malicious inputs.

- *Personalized Cloud Services and Ethical Profiling*

AI-driven personalization enhances service delivery by tailoring experiences to individual user preferences. However, this approach can also reinforce stereotypes and biases, raising ethical concerns about profiling and discrimination. A zero-trust framework, which never assumes user trust, is needed for effective data protection and ethical compliance. [14]. By embedding such frameworks, we can harmonize the benefits of personalization with social issues and ethical concerns.

- *Bias in Decision Logic*

Anomaly detection techniques are popular for detecting security threats and anomalous behavior. However, this does not preclude these models to overestimate minority behaviors and benign activities of marginalized groups making them give unfair and false positives [3]. To provide fairness, fairness constraints and biases mitigation methods can help them in decision logic to help to ensure fairness results and fair AI usage and enforce these AI ethical principles [13]. Ongoing monitoring and auditing are also crucial with fairness over time.

- *Accountability and Governance*

Governance structures are necessary for smart AI deployment in cloud environment. Linthicum promotes alignment of operations, clearly defined roles, and accountability structures to encourage responsible use of AI [1]. Moreover, human-in-the-loop validation procedures can support oversight, error correction and context-sensitive decision-making processes, thus building trust and transparency in automated structures [8].

- *Privacy and Consent in Automated Decisions*

AI system typically uses behavioral and personal characteristics to guide decision making. It is critical to maintain user privacy and to obtain informed consent. Omar stresses the importance of strong privacy protections such as differential privacy, and secure data-handling protocols in the cybersecurity-contextual frameworks in cybersecurity as well as protecting user and to comply with the compliance of data protection regulations and the need to ensure privacy rights [8].

III. CHALLENGES AND ETHICAL CONSIDERATIONS

AI implementation in cloud platforms brings a range of potential ethical pitfalls—from how things are developed to how security activities and automated decision-making are handled. These challenges are not separate -- they cut across technical layers and require orchestrated governance, transparency, and human oversight.

➤ *Bias in Code and Decision Logic*

AI-generated code and decision systems often reflect the biases embedded in their training data. In development, this can manifest as discriminatory logic or exclusionary defaults:

Hypothetical biased suggestion

```
def classify(data): if data['race'] == 'minority':
    return 'High risk'
    else:
    return 'Low risk'
```

In security, skewed anomaly detection can over-identify some user behaviors as threats [2]. In decision-making, personalization algorithms can reinforce stereotypes or create filter bubbles [1]. Such risks need bias audits, fairness constraints, and diverse training sets.

➤ *Opacity and Explainability*

AI systems are often black boxes, making it hard to see how decisions are made. In development, they are limited in debugging and validation. Such transparency in threat categorizations is especially useful in security where false positives or missed attacks can occur due to unarticulated threat

classifications, leading to significant risks [2]. In the context of decision making, when explainability is absent, it leads to lack of trust and regulatory issues on the part of the users [13]. To overcome this divide, techniques such as explainable AI (XAI), model introspection, and natural language rationales are necessary.

➤ *Accountability and Governance*

When AI-generated code causes harm — such as a data breach or service outage — responsibility is often unclear. Consider:

Insecure endpoint example

```
@app.route('/data', methods=['GET'])
def get_data():
    # No authentication
    return jsonify(fetch_sensitive_data())
```

In security scenarios, automated incident response may create a containment environment by initiating an action which interrupts services [8]. During decision-making in resource allocation or access control, errors could cause cascading effects [3]. Ethical governance includes clear role definitions, escalation paths, and human-in-the-loop validation [1]. This includes setting roles, escalation pathways, and human-in-the-

loop governance to be implemented as part of a human-in-the-loop process.

➤ *Security Vulnerabilities from AI Suggestions*

AI-generated code may introduce subtle but critical security flaws. For example:

SQL injection vulnerability

```
def process_transaction(user_id, amount, auth_token):
    if authenticate(user_id, auth_token):
        query = f"UPDATE accounts SET balance = balance + {amount} WHERE user_id = {user_id}"
        cursor.execute(query)
        return True
    return False
```

These examples highlight the need for secure defaults, automated code scanning, and ethical review pipelines.

➤ *Consent and Data Governance*

AI systems depend on user data — access logs, behavioral telemetry, and usage patterns. In development, this data may be utilized to fine-tune models. In security, it drives anomaly detection and identity scoring. In decision-making, it fuels personalization and policy enforcement. Without strong consent mechanisms, these can constitute practices that may potentially break user expectations and violate legal standards. Federated learning methods, differential privacy methods, and consent-aware models are becoming more and more important [2].

➤ *Over-Automation and Human Displacement*

The efficiency of AI can enable over-automation to the detriment of human judgment in critical workflows. In development, then, this tends to lower the scrutiny of code review. In security, it can result in blind faith in automated remediation. The above results may lead to ethical deliberation being marginalized in the decision-making process in favor of algorithmic optimization [8][12].

Ethical deployment requires balancing automation with human oversight, especially in high-stakes environments.

IV. CONCLUSION

Artificial Intelligence is revolutionizing cloud platforms in development, security, and decision-making. Whether in code generation or automation of infrastructure, or threat detection and policy enforcement, AI provides an unparalleled efficiency and scalability. But these abilities also involve some ethical trade-offs that must be weighed in their adoption.

Bias in training data, opacity in decision logic, and accountability gaps are real threats to fairness, safety, and trust. Code created by AI may create vulnerabilities; automated choices may perpetuate discrimination; and security systems can fail to be human controlled. The issues raised are about more than “technical” issues — they are of a deep ethical nature, with the need for transparency in governance, explainability of models in use, and inclusiveness.

➤ *To Ensure Responsible AI in Cloud Platforms, Organizations Must Adopt A Multi-Layered Approach:*

- Audit and validate AI-generated logic before deployment
- Embed fairness and privacy constraints into model design
- Maintain human-in-the-loop oversight for critical decisions
- Establish clear accountability frameworks across stakeholders

At the end of the day, ethical AI isn't an add-on but it's a foundation. Yet, just as cloud platforms need to evolve, so should our devotion to develop systems that are not merely intelligent, but just, secure, and accountable.

As AI continues to shape cloud platforms, ethical governance must evolve in parallel to ensure these systems remain not only intelligent, but also just and trustworthy.

REFERENCES

- [1]. R. Vayyala, “Ethical AI and Analytics in Cloud-Based Data Ecosystems,” *2025 6th International Conference on Artificial Intelligence, Robotics, and Control (AIRC)*, IEEE, pp. 264–267, DOI: 10.1109/AIRC64931.2025.11077557
- [2]. S. Surya, K. Onapakala, D. Santhakumar, V. B. T. Raaj, A. Tyagi, and N. K. Kumar, “AI-Driven Threat Detection: Implementing Multi-Layer Security Networks in Cloud Environments,” *2025 International Conference on Pervasive Computational Technologies (ICPCT)*, IEEE, DOI: 10.1109/ICPCT64145.2025.10941038
- [3]. K. A. Singh and A. Choudhry, “AI-Powered Strategies for Cloud Infrastructure Management,” *2025 4th OPJU International Technology Conference (OTCON)*, IEEE, DOI: 10.1109/OTCON65728.2025.11070393.
- [4]. D. S. Linthicum, “Making Sense of AI in Public Clouds,” *IEEE Cloud Computing*, vol. 4, no. 6, pp. 70–72, Nov./Dec. 2017.
- [5]. A. Polamarasetti, V. Yammanur, N. Ravuri, R. Vadisetty, and V. V. R. Murthy, “Enhancing Cloud Performance with AI-Based Predictive Analytics,” *2025 International Conference on Networks and Cryptology (NETCRYPT)*, IEEE, DOI: 10.1109/NETCRYPT65877.2025.11102149.
- [6]. “Zero-Day Attack Detection Using Machine Learning,” *IEEE Xplore*, 2024. [Online]. Available: <https://ieeexplore.ieee.org/document/11166929/>
- [7]. Santhosh Chitraju Gopal Varma, “AI-Enhanced Cloud Security: Proactive Threat Detection and Response Mechanisms”, IJFMR, Sep. 2024.
- [8]. Marwan Omar, “Integrative Approaches in Cybersecurity, Artificial Intelligence, and Data Management: A Comprehensive Review and Analysis”, Illinois Institute of Technology, 2024. [Online]. Available: <http://arxiv.org/pdf/2408.05888.pdf>
- [9]. Kush Janani, “The Human-Machine Identity Blur: A Unified Framework for Cybersecurity Risk Management in 2025,” *arXiv*, 2025. [Online]. Available: <https://arxiv.org/pdf/2503.18255.pdf>
- [10]. Qian Cheng, Doyen Sahoo, “AI for IT Operations (AIOps) on Cloud Platforms: Reviews, Opportunities and Challenges” *IEEE Xplore & arXiv*, 2023. [Online]. Available: <https://arxiv.org/pdf/2304.04661>
- [11]. Animesh Kumar, “AI-Driven Innovations in Modern Cloud Computing ” *arXiv*, 2024. [Online]. Available: <https://arxiv.org/pdf/2410.15960>
- [12]. Jesu Narkarunai Arasu Malaiyappan; Sanjeev Prakash, “Enhancing Cloud Compliance: A Machine Learning Approach ,” *AIJMR*, 2024. [Online]. Available: <https://www.aijmr.com/papers/2024/2/1036.pdf>
- [13]. “Ethical AI in Cloud Computing: AWS Implementation and Societal Implications,” *International Journal of Computer Engineering & Technology*, vol. 15, no. 6, pp. 110–118, 2024. [Online]. Available: https://iaeme.com/MasterAdmin/Journal_uploads/IJCET/VOLUME_15_ISSUE_6/IJCET_15_06_110.pdf
- [14]. Muhammad Liman Gambo; Ahmad Almulhem, “Zero Trust Architecture: A Systematic Literature Review”, 2025, Cryptography and Security, <https://arxiv.org/abs/2503.11659>