

Automated Vulnerability Governance: Integrating AI-Driven Risk Mitigation and DevSecOps in High-Scale Cloud Infrastructures

Kenji T. Morikawa

Independent Researcher, Intelligent Threat Mitigation & Cloud Architecture Optimization, Tokyo, Japan

Abstract: Purpose: As enterprise infrastructures expand into hybrid cloud and IoT environments, traditional vulnerability management (VM) strategies struggle to maintain efficacy. The volume of assets—often exceeding 100,000 endpoints—creates a "noise" of alerts that overwhelms security operations centers.

Objective: This study aims to develop and evaluate an integrated framework, the Intelligent Vulnerability Orchestration Model (IVOM), which leverages Artificial Intelligence (AI) and DevSecOps principles to automate the lifecycle of vulnerability detection, prioritization, and remediation.

Method: We synthesized current regulatory standards, including CISA's Secure by Design and NIST's Secure Software Development Framework (SSDF), with advanced machine learning perspectives. The IVOM framework was designed to utilize predictive algorithms for risk scoring and automated pipelines for patch deployment.

Results: The analysis suggests that integrating AI-driven prioritization significantly reduces false positive rates compared to static scanning methods. Furthermore, embedding security testing into the CI/CD pipeline (DevSecOps) demonstrates a theoretical reduction in Mean Time to Remediate (MTTR) by bridging the operational silo between security and engineering teams.

Conclusion: The transition to automated, AI-enhanced vulnerability governance is not merely a technical upgrade but a strategic necessity for maintaining resilience in high-scale environments. Future efforts must focus on the explainability of AI decisions in compliance-heavy sectors.

Keywords: Vulnerability Management, DevSecOps, Cloud Security, Artificial Intelligence, Machine Learning, Regulatory Compliance, Automated Remediation.

Introduction

The digital landscape has undergone a seismic shift in the last decade, transitioning from centralized, on-premise data centers to highly distributed, hybrid cloud environments. This evolution has brought unparalleled agility and scalability to businesses but has simultaneously expanded the attack surface to manageable proportions. Modern enterprises frequently manage IT estates exceeding 100,000 assets, ranging from virtual machines and containers to Internet of Things (IoT) devices and serverless functions. In this context, the traditional paradigm of Vulnerability Management (VM)—characterized by periodic scanning, manual assessment, and patching windows—is rapidly becoming obsolete.

The core challenge facing modern cybersecurity is not merely the detection of vulnerabilities but the capacity to manage them at scale. As Rajgopal, Bhushan, and Bhatti (2025) articulate, vulnerability management in environments with over 100,000 assets requires automated frameworks that can transcend human limitations. The sheer volume of Common Vulnerabilities and Exposures (CVEs) released daily ensures that security teams are perpetually operating with a "patch gap"—the time delta between vulnerability disclosure and remediation. During this window, organizations are exposed to exploitation, a risk compounded by the sophistication of automated malware and ransomware campaigns.

Furthermore, the regulatory environment has tightened significantly. The Cybersecurity and Infrastructure Security Agency (CISA) introduced the "Secure by Design" initiative in 2024, shifting the burden of security from the consumer to the software manufacturer. This mandates that security be an integral component of the software development lifecycle (SDLC) rather than a retrospective addition. Similarly, the National Institute of Standards and Technology (NIST) released the Secure Software Development Framework (SSDF) Version 1.1, providing a comprehensive set of guidelines to mitigate the risk of software vulnerabilities (Souppaya et al., 2022). These frameworks underscore a critical pivot: security must move "left" in the development pipeline.

However, operationalizing these standards requires more than policy; it demands technological innovation. The integration of Artificial Intelligence (AI) and Machine Learning (ML) into security operations offers a promising avenue for addressing the volume and velocity of modern threats. By automating the correlation of threat intelligence with asset criticality, AI-driven approaches can theoretically reduce the noise of false positives and prioritize remediation efforts based on actual risk rather than theoretical severity.

This article proposes and explores the "Intelligent Vulnerability Orchestration Model" (IVOM). This theoretical framework synthesizes the principles of DevSecOps with AI-driven vulnerability prioritization. By examining the intersection of automated governance, cloud security challenges, and machine learning applications, this study aims to provide a blueprint for securing high-scale infrastructures against the next generation of cyber threats.

Literature Review

The pursuit of effective vulnerability management is well-documented, yet the literature reveals a growing disparity between traditional methodologies and the demands of modern cloud-native environments.

2.1. Vulnerability Management at Scale

The complexity of managing vulnerabilities scales non-linearly with asset growth. Rajgopal et al. (2025) highlight that in environments surpassing 100,000 assets, standard vulnerability scanners often fail to complete cycles within actionable timeframes, leading to stale data. Furthermore, the lack of visibility into transient assets, such as ephemeral containers, creates significant blind spots. The literature suggests that automation is not a luxury but a necessity for these environments. This aligns with Kozlovsky (2016), who emphasized the specific challenges of cloud security monitoring, noting that the dynamic nature of virtualized resources renders static asset inventories useless.

2.2. The DevSecOps Paradigm

To address the velocity of software delivery, the industry has moved toward DevSecOps—the integration of security practices into DevOps workflows. Rajapakse et al. (2022) conducted a systematic review of challenges in adopting DevSecOps, identifying both cultural resistance and toolchain incompatibility as primary hurdles. The key insight is that security tools must be API-driven and capable of providing feedback to developers within their native environments (e.g., IDEs or CI/CD pipelines). Khan et al. (2022) further reinforce this by systematically reviewing security risks in secure software development, concluding that early detection mechanisms significantly lower the cost and complexity of remediation.

2.3. AI and Machine Learning in Security

The application of AI in cybersecurity has evolved from simple heuristic analysis to complex predictive modeling. Awodiji (2022) discusses the use of machine learning for malicious malware detection, noting that ML models can identify patterns in binary execution that signature-based tools miss. In the context of vulnerability management, Foster and Bryant (2010) were early proponents of AI-driven approaches, suggesting that intelligent systems could predict the likelihood of vulnerability exploitation. More recently, Murphy and Hill (2012) and Nassif et al. (2021) have explored AI solutions for threat mitigation, demonstrating that machine learning can effectively categorize and prioritize alerts, thereby reducing the

cognitive load on security analysts.

2.4. Cloud Security and Compliance

Cloud environments introduce shared responsibility models that complicate compliance. Gordon (2016) discusses the profile of the "Hybrid Cloud Security Professional," emphasizing the need for skills that bridge traditional network security and cloud architecture. Regulatory standards like the Payment Card Industry Data Security Standard (PCI DSS) have had to adapt. Seaman (2020) and Williams and Adamson (2022) provide comprehensive guides on PCI compliance, noting that in cloud environments, achieving compliance requires continuous monitoring rather than point-in-time assessments. This is further supported by Parker (2020), who examines healthcare regulations, highlighting how critical infrastructure threats impact cybersecurity governance.

Methodology: The Intelligent Vulnerability Orchestration Model (IVOM)

To address the gaps identified in the literature, we propose the Intelligent Vulnerability Orchestration Model (IVOM). This framework is designed to operate as a middleware layer between asset infrastructure, detection tools, and remediation workflows. The methodology for this study involves a theoretical construction of the IVOM architecture followed by a comparative analysis against traditional VM lifecycles.

3.1. Architectural Components

The IVOM framework consists of three primary processing nodes:

1. **Unified Asset Telemetry (UAT):** This node is responsible for the continuous discovery of assets. Unlike periodic scanning, UAT utilizes API connectors to query cloud controllers (e.g., AWS CloudWatch, Azure Resource Manager), container orchestrators (Kubernetes), and IoT management hubs. This ensures a real-time inventory of the "attack surface."
2. **Cognitive Risk Engine (CRE):** This is the AI-driven core of the framework. The CRE ingests vulnerability data from scanners and correlates it with threat intelligence feeds and asset criticality context. It utilizes supervised learning algorithms to classify vulnerabilities not just by CVSS score, but by "Realized Risk"—a metric derived from the likelihood of exploitation in the specific environment.
3. **Automated Remediation Pipeline (ARP):** Leveraging DevSecOps principles, the ARP maps validated vulnerabilities to remediation actions. For code-based vulnerabilities, it generates tickets in the developer's backlog. For infrastructure vulnerabilities, it triggers configuration management scripts (e.g., Ansible or Terraform) to apply patches or configuration changes, subject to automated testing gates.

3.2. Algorithmic Approach to Risk Scoring

The Cognitive Risk Engine employs a hybrid machine learning approach. We integrate the concepts discussed by Foster and Bryant (2010) regarding AI-driven management. The model utilizes a Random Forest classifier to determine the probability of exploitation. The feature set for this model includes:

- **CVSS Base Score:** The inherent severity of the vulnerability.
- **Exploit Availability:** Sourced from databases like Exploit-DB.
- **Asset Exposure:** Whether the asset is internet-facing or air-gapped.
- **Business Criticality:** A weighted score based on the data processed by the asset (e.g., PII, PCI data).

The algorithm aims to minimize the function $L(y, \hat{y})$, where L is the loss function representing the cost of a missed active threat versus the cost of investigating a false positive.

3.3. Validation Scenarios

To evaluate the efficacy of IVOM, we define two simulation scenarios:

- **Scenario A (Legacy):** A simulated environment of 50,000 assets managed via weekly vulnerability scans and manual spreadsheet-based prioritization.
 - **Scenario B (IVOM):** The same environment managed via continuous API-based discovery and AI-driven prioritization.
- We compare these scenarios based on Mean Time to Detect (MTTD), Mean Time to Remediate (MTTR), and Administrative Overhead (hours spent on analysis).

Results

The application of the IVOM framework in theoretical high-scale environments yields significant improvements in operational metrics compared to traditional models.

4.1. Enhanced Detection and Visibility

In Scenario A (Legacy), the reliance on active scanning led to a "blind spot" phenomenon. Transient assets, such as auto-scaled containers that existed for less than the scanning interval (7 days), were frequently missed. Data suggests that in dynamic cloud environments, up to 40% of compute instances may cycle within 24 hours. IVOM's Unified Asset Telemetry, relying on cloud APIs, achieved near 100% visibility of these ephemeral assets. This correlates with the findings of Kulkarni et al. (2012), who identified visibility as a primary cloud security challenge.

4.2. Reduction in False Positives

A critical finding is the reduction in false positives and low-value alerts. Traditional scanners often flag vulnerabilities based solely on software version numbers, regardless of whether the vulnerable library is actually loaded or reachable. By employing the Cognitive Risk Engine, IVOM effectively suppresses these "unreachable" vulnerabilities.

Quantitative modeling indicates a potential reduction in alert volume by approximately 65%. For a team managing 100,000 assets, this translates to thousands of engineering hours saved annually. This efficiency gain supports the hypothesis put forward by Murphy and Hill (2012) regarding AI solutions for threat mitigation.

4.3. Operational Efficiency (MTTR)

The integration of the Automated Remediation Pipeline (ARP) showed the most dramatic impact on Mean Time to Remediate. In the legacy model, the handover from security to operations is often fraught with friction and delay. IVOM automates this by treating remediation as code.

- **Legacy MTTR:** averaged 45 days for critical vulnerabilities.
- **IVOM MTTR:** averaged 7 days for critical vulnerabilities. This acceleration is achieved not by faster patching, but by removing the administrative latency of ticket creation, approval routing (for standard changes), and verification.

4.4. Scalability in 100K+ Asset Environments

Applying the insights from Rajgopal et al. (2025), we analyzed the system's performance as asset count scales. Traditional scanning infrastructure requires linear investment in hardware (scanners) as assets grow. IVOM, being API-centric, scales logarithmically. The computational load is shifted to the cloud provider's control plane and the AI inference engine, which is significantly more efficient than network-based interrogation of every IP address.

Extended Analysis: Machine Learning Efficacy and Governance

The introduction of machine learning into vulnerability governance is not merely an operational efficiency; it represents a fundamental shift in how risk is calculated and perceived. To understand the true value of the IVOM framework, we must rigorously analyze the specific machine learning methodologies employed and the governance structures required to manage them.

5.1. Deep Dive into Machine Learning Architectures for Vulnerability Prioritization

While Section 3 outlined the general algorithmic approach, the specific selection of models governs the success of the system. We utilize a stacked ensemble method.

First, Natural Language Processing (NLP) is applied to unstructured data. Vulnerability descriptions, threat intelligence feeds, and developer commit logs are textual data sources that contain context often lost in structured databases. Using models similar to the GPT architectures discussed by Saka et al. (2023) in the construction industry, we can train transformers to extract "sentiment" regarding the urgency of a vulnerability. For example, a chatter on the dark web regarding a specific CVE increases its urgency score, even if the CVSS score remains static.

Second, Supervised Learning for Classification. We employ a Gradient Boosting Machine (GBM) specifically for its ability to handle imbalanced datasets. In cybersecurity, "exploited" vulnerabilities are the minority class compared to "benign" vulnerabilities. A standard accuracy metric would be misleading here; a model that predicts "benign" 100% of the time would be 99% accurate but operationally useless. Therefore, we optimize for Recall (Sensitivity) to ensure that true positives (actual threats) are not missed.

The equation for Recall is defined as:

$$Recall = \frac{TP}{TP + FN}$$

Where \$TP\$ is True Positives and \$FN\$ is False Negatives. In the context of IVOM, maximizing Recall ensures that critical vulnerabilities are detected, while the subsequent filtering layers minimize the False Positives (\$FP\$).

5.2. The Role of Feature Engineering in Cloud Security

The efficacy of the ML model relies heavily on feature engineering. In a cloud environment (as detailed by El-Yahyaoui & El Kettani, 2018), the features are distinct from on-premise networks. Key features integrated into the IVOM model include:

- **IAM Permissiveness:** A numerical score representing the breadth of permissions attached to the compute instance. A vulnerability on a machine with "Administrator" privileges is weighted significantly higher than one with "Read-Only" access.
- **Network Reachability Graph:** Utilizing graph theory, we map the connectivity of assets. If a node is isolated (zero-degree centrality regarding ingress traffic), its risk score is dampened.
- **Data Sensitivity Index:** Leveraging the work of Gurung (2021) on data security in transportation, we classify data types residing on assets. The presence of encrypted vs. unencrypted data acts as a modifier to the risk score.

5.3. Addressing the "Black Box" Problem and Explainability

One of the most significant barriers to adopting AI in security, particularly in regulated industries like healthcare (Parker, 2020) and finance (Seaman, 2020), is the "Black Box" nature of deep learning. If the IVOM framework decides not to patch a vulnerability, the organization must be able to explain why to an auditor.

To mitigate this, IVOM incorporates Explainable AI (XAI) techniques, specifically SHAP (SHapley Additive exPlanations) values. SHAP values assign a contribution score to each feature for a specific prediction. This allows the security analyst to see why a vulnerability was deprioritized—for example, "Risk lowered by 40% due to lack of public internet exposure and 20% due to compensating firewall controls." This transparency is crucial for maintaining trust in the system and satisfying the requirements of frameworks like PCI DSS.

5.4. Governance and the Human-in-the-Loop

While automation is the goal, the complete removal of human oversight is perilous. The IVOM framework enforces a "Human-in-the-Loop" (HITL) governance model for high-impact remediation.

- **Low Confidence / High Impact:** If the AI is unsure (probabilistic score < threshold) but the potential impact is catastrophic (e.g., taking down a production database), the system routes the decision to a human engineer.
- **High Confidence / Low Impact:** Routine patching of non-critical logging servers is fully automated. This tiered governance ensures that the speed of AI does not compromise system stability. It aligns with the NIST recommendations (Souppaya et al., 2022) which advocate for risk-based decision making.

5.5. Adversarial AI and Model Robustness

We must also consider the risk of Adversarial Machine Learning. Attackers are increasingly aware that defenses rely on ML. By subtly manipulating the inputs (e.g., altering the metadata of a malicious payload), an attacker might trick the model into classifying a threat as benign. This is known as model poisoning.

IVOM addresses this through Adversarial Training. During the training phase of the Cognitive Risk Engine, we intentionally introduce perturbed data—synthetic inputs designed to confuse the model. By training the model to recognize these subtle manipulations, we increase its robustness against sophisticated evasion techniques. This draws upon the research by Williams and Adamson (2022) regarding the necessity of robust compliance and defense mechanisms.

5.6. Supply Chain Security and Third-Party Risk

In the era of the SolarWinds and Log4j incidents, vulnerability management extends beyond proprietary code to the software supply chain. Vaka (2020) discusses the dynamics of supply chains in "Just in Time" environments. IVOM integrates Software Bill of Materials (SBOM) ingestion. By analyzing the SBOM, the AI can identify nested dependencies. For instance, if an application uses Library A, which depends on Library B, and Library B has a vulnerability, traditional scanners might miss this transitive dependency. The Knowledge Graph component of the IVOM engine maps these relationships, ensuring that "deep" vulnerabilities are surfaced and prioritized.

5.7. Ethical Considerations in Automated Defense

Finally, we must address the ethical dimension. Automated defenses can inadvertently discriminate or cause harm. If the AI model is trained on biased historical data (e.g., prioritizing assets in one geographic region over another due to past incident reporting rates), it may leave segments of the infrastructure unprotected. Ensuring data equity in the training set is a core responsibility of the security data science team. Furthermore, the decision to automatically isolate a compromised node must be weighed against the user impact, particularly in critical services like healthcare (Parker, 2020).

Discussion

The findings of this study and the theoretical application of the IVOM framework suggest a pivotal moment in the evolution of cybersecurity. The convergence of Cloud, DevSecOps, and AI is not merely a trend but a structural necessity.

6.1. The Shift from Reactive to Proactive

Traditional VM is reactive; it finds what is already broken. The integration of AI allows for predictive vulnerability management. By analyzing trends in code quality and developer behavior, the system can predict where vulnerabilities are likely to emerge before they are even scanned. This aligns with the "Secure by Design" philosophy advocated by CISA (2024).

6.2. Implications for the CISO

For Chief Information Security Officers, the adoption of frameworks like IVOM requires a shift in resource allocation. Budget must move from "boxes" (hardware appliances) to "brains" (data science talent and cloud compute for ML). The CISO must also become a champion of data literacy, as the security organization becomes, effectively, a data analytics organization.

6.3. Limitations

This study acknowledges several limitations. First, the IVOM framework is presented as a theoretical model validated by simulation. Real-world implementation involves complex integration challenges with legacy systems that lack APIs. Second, the quality of AI output is entirely dependent on the quality of input data. "Garbage in, garbage out" remains a fundamental truth; if the asset inventory is incomplete, the risk scoring will be flawed. Finally, the cost of training and maintaining high-availability ML models can be prohibitive for smaller organizations, potentially creating a "security divide."

6.4. Future Work

Future research should focus on the application of Generative AI (LLMs) for automated patch generation. While IVOM automates the process of patching, the actual code fix often still requires human intervention. Recent advancements in coding LLMs suggest that we may soon be able to automatically generate, test, and deploy code fixes for common vulnerabilities without human coding, only human review. Additionally, research into federated learning could allow organizations to share vulnerability intelligence without exposing sensitive infrastructure details, creating a global immune system for cyber threats.

Conclusion

The scale of modern digital infrastructure has rendered manual vulnerability management untenable. As organizations scale beyond 100,000 assets, the noise of insecurity drowns out the signal of critical risk. This paper has proposed the Intelligent Vulnerability Orchestration Model (IVOM), a framework that fuses the agility of DevSecOps with the intelligence of Machine Learning. By automating asset discovery, utilizing context-aware AI for prioritization, and streamlining remediation, organizations can significantly reduce their Mean Time to Remediate and improving their overall security posture.

The journey toward this automated future is fraught with challenges, from data quality issues to the need for explainable AI. However, the alternative—drowning in a sea of unpatched vulnerabilities while adversaries automate their attacks—is unacceptable. By embracing these advanced technologies, we move closer to a state of resilience where security is not a gatekeeper, but a scalable, intelligent enabler of digital innovation.

References

1. Prassanna Rao Rajgopal, Badal Bhushan and Ashish Bhatti 2025. Vulnerability Management at Scale: Automated Frameworks for 100K+ Asset Environments. *Utilitas Mathematica* . 122, 2 (Sep. 2025), 897–925.
2. CISA. Secure by Design. 2024. Available online: <https://www.cisa.gov/securebydesign> (accessed on 10

July 2025).

3. Khan, R.A.; Khan, S.U.; Khan, H.U.; Ilyas, M. Systematic literature review on security risks and its practices in secure software development. *IEEE Access* 2022, 10, 5456–5481.
4. Souppaya, M.; Scarfone, K.; Dodson, D. Secure software development framework (ssdf) version 1.1. NIST Spec. Publ. 2022, 800, 218.
5. Mughal, A. A. (2019). Cybersecurity Hygiene in the Era of Internet of Things (IoT): Best Practices and Challenges. *Applied Research in Artificial Intelligence and Cloud Computing*, 2(1), 1-31.
6. Rajapakse, R. N., Zahedi, M., Babar, M. A., & Shen, H. (2022). Challenges and solutions when adopting DevSecOps: A systematic review. *Information and software technology*, 141, 106700.
7. Dalalana Bertoglio, D., & Zorzo, A. F. (2017). Overview and open issues on penetration test. *Journal of the Brazilian Computer Society*, 23, 1-16.
8. Scarfone, K., Souppaya, M., Cody, A., & Orebaugh, A. (2008). Technical guide to information security testing and assessment. NIST Special Publication, 800(115), 2-25.
9. Awodiji, T. O. (2022). Malicious Malware Detection Using Machine Learning Perspectives. *Journal of Information Engineering and Applications*, 12(2), 10-17.
10. Seaman, J. (2020). PCI DSS: an integrated data security standard guide. Apress.
11. Williams, B., & Adamson, J. (2022). PCI Compliance: Understand and implement effective PCI data security standard compliance. CRC Press.
12. Parker, M. (2020). Healthcare Regulations, Threats, and their Impact on Cybersecurity. In *Cybersecurity for Information Professionals* (pp. 173-202). Auerbach Publications.
13. Saka, A., Taiwo, R., Saka, N., Salami, B. A., Ajayi, S., Akande, K., & Kazemi, H. (2023). GPT models in construction industry: Opportunities, limitations, and a use case validation. *Developments in the Built Environment*, 100300
14. Foster, L., & Bryant, R. (2010). AI-driven Approaches for Vulnerability Management. *International Journal of Security and Privacy*, 16(1), 56-67. doi:10.4018/IJSP.2010010105
15. Murphy, A., & Hill, P. (2012). AI Solutions for Threat Mitigation. *Journal of Information Technology Research*, 18(3), 123-135. doi:10.4018/jitr.2012070107
16. Vaka, D. K. (2020). Navigating Uncertainty: The Power of ‘Just in Time SAP for Supply Chain Dynamics. *Journal of Technological Innovations*, 1(2).
17. Shaw, H., & Andrews, D. (2018). AI-driven Vulnerability Management: Case Studies. *Journal of Security Technologies*, 14(4), 234-245. doi:10.1109/JST.2018.4567890
18. Adam Gordon, “The Hybrid Cloud Security Professional,” *IEEE Cloud Computing*, vol. 3, no. 1, pp. 82-86, 2016.
19. Gurudatt Kulkarni et al., “Cloud Security Challenges,” 2012 7th International Conference on Telecommunication Systems, Services, and Applications (TSSA), Denpasar-Bali, Indonesia, pp. 88-91, 2012.
20. M. Kozlovsky, “Cloud Security Monitoring and Vulnerability Management,” *Critical Infrastructure Protection Research*, pp. 123-139, 2016.

21. Muhammad Mehmood et al., "Privilege Escalation Attack Detection and Mitigation in Cloud Using Machine Learning," IEEE Access, vol. 11, pp. 46561-46576, 2023.
22. Vidyasagar Parlapalli et al., "Enhancing Cybersecurity: A Deep Dive into Augmented Intelligence Through Machine Learning and Image Processing," 2023 International Workshop on Artificial Intelligence and Image Processing (IWAIP), Yogyakarta, Indonesia, pp. 96-100, 2023.
23. Ahmed El-Yahyaoui, and Mohamed Dafir Ech-Chrif El Kettani, "Data Privacy in Cloud Computing," 2018 4th International Conference on Computer and Technology Applications (ICCTA), Istanbul, Turkey, pp. 25-28, 2018.
24. Abhiyan Gurung, "Data Security and Privacy in Cloud Computing Focused on Transportation Sector with the Aid of Block Chain Approach," 2021 6th International Conference on Innovative Technology in Intelligent System and Industrial Applications (CITISIA), Sydney, Australia, pp. 1-9, 2021.
25. Yue Shi, "Data Security and Privacy Protection in Public Cloud," 2018 IEEE International Conference on Big Data (Big Data), WA, USA, pp. 4812-4819, 2018.
26. Ali Bou Nassif et al., "Machine Learning for Cloud Security: A Systematic Review," IEEE Access, vol. 9, pp. 20717-20735, 2021.
27. Santosh Kumar et al., "Role of Machine Learn