

Srinivasa Rao Maka¹,
 Suneel Babu Boppana,
 Gangadhar Sadaram,
 Niharika Katnapally,
 Laxmana Murthy
 Karaka,
 Manikanth Sakuru

Automating Cyber Threat Response Using Agentic AI and Reinforcement Learning Techniques



Abstract: The essay proposes an agentic artificial intelligence (AI) and reinforcement learning (RL) framework, design and validation approaches to significantly automate response to continued cyber-attacks. First-generation AI tools are analyzed and then variations of them that qualify them as cyber-agents are proposed. A new typology and a 11-dimensional notation framework for them are used. A typology and a 11-dimensional notation framework for nascent, cyber-agentic anti-antagonistic hulls response tree (BRT)-based AI functions are used and the automatic location of them with an alliance of tabu search (TS) techniques is obtained. On the outcomes derived from the case of the international cyber-attack against the ESB in 2021, the validation of the research is presented, along with potential directions for future work .

An escalating sequence of cyber-attacks have been waged against Ireland and the ESB, attempting to overwhelm the relatively nascent national cyber-security infrastructures. Consequently, an intensified research effort in the country to find new means and tools to enhance its national cybersecurity measures is under way. It would be useful to automate emergency response to cyber-attacks. As a consequence of evolving AI research, it became possible to propose variations of cyber-defender architectures that qualify them as cyber-agents; the simpler bots can be cyber-agents. Up to now there is no comprehensive taxonomy, typology, or notation framework that correctly accounts for the new cyber-agentic AI. A new typology and an 11-dimensional notation framework for them are proposed. Intense work is under way to make it possible to automatically design those cyber-agents. Another typical case is the USA. The American federal government, tired of the excesses of a few multinational rogue corporations, gave those companies an ultimatum about their cyber-violations of American law.

Keywords: Agentic AI, artificial intelligence, cybersecurity, cyber threat response, reinforcement learning, security operations, Agentic AI, Cyber Threat Automation, Reinforcement Learning, Threat Detection Algorithms, AI-based Cybersecurity, Autonomous Threat Mitigation, AI in Cyber Defense, Intelligent Cyber Response, Self-learning Security Systems, Adaptive Security Frameworks.

1. Introduction

Huge advancements in information technology over the past few decades have resulted in a doubling of cyber attacks yearly. To protect against these threats, governments, financial institutions, and commercial agencies are forced to allocate billions annually into developing sophisticated security measures and an entire cyber defense industry. Nevertheless, cyber threats of varying complexity, whether state-sponsored, criminal, or by hacktivist

¹ 1., North Star Group Inc, Software Engineer

2., iSite Technologies, Project Manager

3., Bank of America, DevOps/ OpenShift Admin Engineer

4., Amazon AWS, BI Developer

5., Microsoft, Senior Support Engineer

6., JP Morgan Chase, Lead Software Engineer

groups, continue to bypass even the most advanced systems, often leading to disastrous consequences. The field of artificial intelligence provides a strong foundation for creating an innovative and vastly superior approach to combat cyber threats, using capable autonomous agents to develop agentic intelligent cyber defense networks. Automating cyber threat response action has been espoused for some time. Traditionally, rule-based reactive mechanisms have been employed in wiping or isolating infected systems and blacklisting malicious entities. Currently, AI techniques are bundling rules and heuristics in combating cyber threats. The realm of Artificial Intelligence in cybersecurity extends far wider than heuristic and rule-based processes.

Amid a vast well-established body of research on cybersecurity and on the application of AI to cybersecurity, the literature on agentic AI approaches to cybersecurity is disjointed and partial. Yet the tension and interplay between the design and brokering of agents in peer-to-peer networks and the cyber threat responses of those agents raise important questions, both for the development of more capable intelligent cyber defense networks and for the prevention of possible negative consequences of their deployment. In addressing this neglect, the specific research objectives and questions of the essay are outlined. Policymakers, commercial defenders, and academics are increasingly interested in perfecting the realm of reinforcement learning techniques due to the emergent and unknown nature of cyber attacks. This is not least because of the almost unlimited capacity for generating hypotheses and optimal strategies in a fully modeled cyber environment. However, in observation replicating closely the digital twin of real-time settings, defender agencies are only just beginning to become aware of and seek to shape or react to reinforcement-based threat responses by cyber attackers. Consequently, this essay seeks to contribute in bridging an important gap by exploring how reinforcement techniques (RT) and learning are being employed in cyber-offense, the defensive implications and challenges thereby raised, and how these findings can be leveraged broader for the design of better ICT and organizational tools and practices in cyber defense.



Fig 1: AI Agents in Cybersecurity (AI SOC)

2. Background and Literature Review

With the evolution and expansion of technology, the dependence on the cyber world is constantly growing, thus leading to increase in potential cyber threats and greater susceptibility/receptiveness to various types of cyber-attacks. Broadly, cyber threats can be primarily characterized by some of the most common attacks including zero day vulnerabilities, social engineering attacks, ransomware, etc., to list down a few is a challenge in itself; making it more complicated, the malicious intent they are equipped with keeps on adapting and getting worse day after day. Various preventive mechanisms have been adopted from time to time to safeguard the world from the wrath of these cyber predators; however, past records of cyber-attacks and malware outbreaks show that such incidents have been increasing at an alarming rate in terms of volume and sophistication too.

As a counter-response mechanism, automated, intelligent, and self-sufficient reactive systems are sought as a panacea. As an attempt to make this happen, Agentic AI has found its application gracefully in the cyber world to come up with contrasting policies against the malicious agent in the cyberworld. The agentic AI aimed at the cyber world learns the policies or rules that can be put forward to teach that cyber world agent. Because of the nonce nature of the cyber world, the automated policies learnt also must be dynamic and ablative so that the cyber world is under attack against any modern world agential wrath.

As the reimplementing of this principle to the cyber world (cybergenetics) booms, the attacker and defender side's cyber activities have started showing some form of rationality. Thus, agentic AI techniques can be admirably implicated in the cybernetic world to either prove the cybernity or put forward divergent policies against it. A simple flow chart, possesses an overview of the proposed policy formation using agentic AI in reaction to the diverse sophisticated impending threat in the cyber world. Merging this principle with contrasting policies either can end the attack in the cyber world. However, the improbably immense component involved in it

facilitates exploring further in directional approaches, possible hardware mechanisms, automated prediction of agential behavior in both attacker and defender agents, etc.. Subsequently, this will spur extensive research in this field, leading to substantial development in counteracting the cyber attacks in the world.

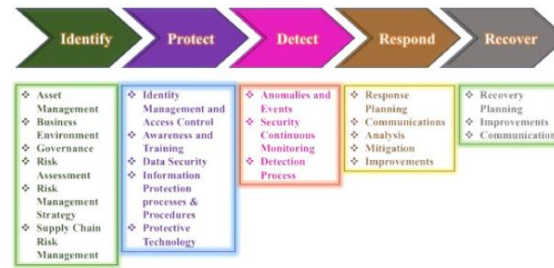


Fig 2: Artificial intelligence for cybersecurity Literature review

2.1. Cyber Threat Landscape

Disruptive technologies and digitalization have become part of the new normative societal behavior. Equally, the cyber threat landscape has become more challenging, as recent cyber-attacks show. It is essential to evaluate the current cyber threat landscape to understand the effectiveness of automating cyber threat response using agentic AI and reinforcement learning techniques. Cyber threats involve a number of different elements, which are described below. The central discussion items are characterized with respect to malware, phishing and denial-of-service (DoS) attacks occurring in the cyber threat landscape. Despite the mentioned categorization, cryptojacking, application layer trouble, advanced persistent threats, insider threats, snap2profit, Web Shell, sem-crude oil pipeline, robot exploits, and heart-breaking can also be categorized under other types of attack, but they are discussed in the context of malware, phishing, and DoS attacks consistent with the mentioned ones.

Colonial Pipeline was targeted by a ransomware attack by a ransomware group, which resulted in the temporary shutdown of their infrastructure and systems. Colonial Pipeline paid the attackers a ransom, which highlights the level of damage that recent vintages of cyber-attacks can have on a company. SolarWind suffered a cyber-attack that injected a backdoor malware in their software, which affected a significant number of companies, which highlights the spread of some malware. Samsung suffered a massive data leak that included source code of devices and claimed that the actor was just an external partner, which highlights insider threats and risks associated with partner companies.

Cyber threats are dynamic and continuously evolve, as they constantly improve in sophistication, as organizations also continue on their defenses, which leads some to argue with some sort of logic that they are still effectively defending themselves on that front.

2.2. Agentic AI in Cybersecurity

As the present technology landscape witnesses a transition from artificial intelligence (AI) as a domain-specific tool to AI-driven systems capable of making decisions autonomously, a changing landscape in cybersecurity is evident. While traditionally focusing on the prevention of attacks by relying on classic Machine Learning and Expert Systems, established components bolster complementary agentic AI-based systems. Specifically, leveraging Reinforcement Learning (RL) components, such systems provide capabilities for adaptive decision making, dynamic planning, task auto-generation and performance assessment. The ensemble orchestrates a machine learning agent capable of consistently monitoring the environment, learning effectively through trial and error and self-improvement by realistic simulation scenarios.

The supervised learning approach utilizing both the development and test datasets is the most popular in the present cybersecurity domain. However, such a conventional element is not adequate as an adaptive and reactive solution responding to rapidly changing cyber threats. The agentic AI-driven solution proposed here focuses specifically on transfer learning, reducing the defensive process adaption time to changing system conditions. The feasibility of such an approach is proved after thorough experimentation on the DARPA 1998 and 1999 datasets. While well adapted to the early scenario of learning and planning, the other conglomerations in a “no-knowledge agent” manner, playing only in the learning environment and providing feedback only on executed tasks, remain

more focused on the extension of exploration strategies to unseen states during attack prediction formation processing .

Agentic AI, enhanced with RL components, may radically transform the cybersecurity domain, which has been under drastic pressure from the rising number of cyber incidents and the increasing sophistication of cyber-attack methods. This paper provides a holistic research contribution capable of grasping the emerging cybersecurity process. As the readiness catalyst, such a systematic approach ensures the full operation of agentic AI-driven systems designed to act against emerging cyber threats effectively.

Equ 1: Multi-Agent Systems (Cooperative or Competitive)

$$R_{\text{total}}(s_t, a_t) = \sum_{i=1}^N R_i(s_t, a_t) \quad \text{Where } N \text{ is the number of agents.}$$

3. Theoretical Framework

The increase in human dependency on information technologies has raised a new set of security concerns associated with protecting the integrity, confidentiality, availability, and authenticity of data. The threat landscape is vast, with numerous actors targeting businesses, individuals, and critical infrastructure, amongst others, and an equal diversity of attack vectors at their disposal. The 2020 SolarWinds supply chain cyber-attack demonstrated both the complexity involved in the execution of sophisticated, multi-stage attacks and the consequences when it is successful. To this extent, AI-based cyberspace defense approaches have been investigated.

Cybersecurity is essential, especially as the reliance on digital technology continues to grow. To protect critical infrastructure and sensitive data, including protecting secure nations, it is essential to maintain stable cybersecurity. AI and cyber defenses can operate effectively together, with the former able to utilize the advantage of threat identification, network monitoring, and effect response algorithm adaptability and speed unmatched by a human defender. These systems will not replace human network defenders but can support them, augmenting their abilities. The concept of AI agents in cybersecurity is broadly akin to an application acting as a self-functioning cyber defender capable of entering an environment to perform actions and containing multiple behaviors. Whilst this does segregate it from the host cyber defense systems that are generally reactive in deployment, the AIs do adapt to the developing threat environment and reprogram their policy accordingly. Purposely developed RL algorithms power the AI agents to enable said environmental interaction in the form of action selection whilst learning optimal policy from the deterministic temporal difference reward signal received retrospectively on policy action completion. This understanding has prompted many works.

3.1. Agentic AI: Concepts and Applications

Much has been written on how talented employees use their unique skills in a highly-flexible manner to effectively respond to cybersecurity threats. But what if an "employee" would be an agentic AI, which uses learned experience through deep reinforcement learning to assess potential threats and then take an autonomous course of action; irrespective of a team mate's wishes? Driven by the work of researchers, this is both a deeply alarming and an incredibly fascinating question. Agentic AI is a first of its kind méthode developed in an effort to address the issue of embracing widespread AI adoption, while demonstrating its profound implications. It refers to fully autonomous AI systems that utilize learned experiences to assess a situation and make decisions on the course of action to be taken based on that assessment. As a result, seemingly subjective and personalized reasoning becomes both deterministic and traceable. Recent developments in the field of artificial intelligence suggest a radical possibility: the swarming of lifeless algorithms through millions of real-time sensors – after the Internet of Things can serve as a new kind of AI-enhanced nervous system. Such agentic AI systems run without intervention and engage the unpredictable 90 percent of human security experts' responses today. Over the last forty years since artificial intelligence was launched, computer science has demonstrated a variety of such capabilities. Consequently, such systems can engage in systematically gathering unspotted data, and accurate, reliable action. By virtue of it being so intimately related to the broader trends of digitalization in society, the proliferation of AI technology in domains of security and defense is under-regulated, and has far-reaching and profound consequences. However, despite their technical sophistication, these

technologies are easily exploited by state and non-state actors, and their potential manipulative misuse renders them gravely insecure.

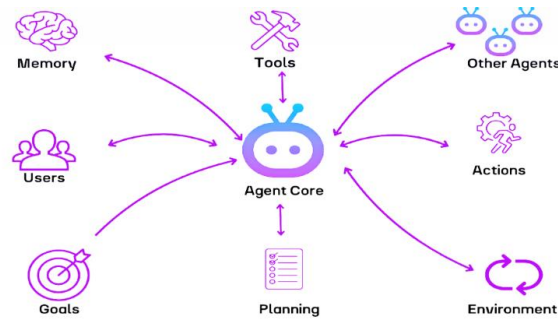


Fig 3: Agentic AI

3.2. Reinforcement Learning: Fundamentals and Applications

Reinforcement

learning is a process by which agents learn optimal behavior in a dynamic environment by taking actions and observing the effects of those actions. The environment is an arena in which the agent operates and gains valuable experiences while receiving rewards in response to the actions taken. Typically, a feedback loop driven by the agent in the environment dictates how the environment receives a state, produces a reward, and sends back the new state as a consequence of the action taken. The agent updates its internal policy on the observed reward signal and the obtained experiences from the executed actions.

There are multiple applications of reinforcement learning that can help in enhancing security protocols or responses. Researchers have proposed different approaches such as using reinforcement learning for optimizing the design of a security protocol, deploying reinforcement learning in combining attack graphs to detect anomalies in the network, and automatically generating responses to cyber threats by employing reinforcement learning for the attacker model. Moreover, reinforcement learning can adapt the security strategy with real-time data, which can be useful in implementing a robust strategy against the external threat. While promising, there are many challenges that researchers currently need to address when deploying reinforcement learning for a security setting such as the convergence of a recommendation algorithm for continuous actions; defending the system from motivated and sophisticated attackers deploying countermeasures against the training signal; ensuring agent stability within the environment; and managing the complexity of the environment, as cybersecurity is too vast and diverse. Yet, reinforcement learning remains a powerful tool that continually shapes and adapts the strategy in the light of the realization of the threat, learning through trial and error under delayed reward.

4. Methodology

This section outlines the research design and approaches utilized in this study to achieve its objectives. To support the analysis, strategies for data collection are detailed, ensuring a robust foundation of evidence. Preprocessing techniques are implemented to cleanse and organize the data prior to modeling, thereby facilitating the cleaning of data and making it ready for use. Details of the model development process are also provided. Algorithms for collective AI and reinforcement learning are selected to be implemented. Descriptions of the training techniques are the focus, suitable for the purpose of enhancing model performance and adaptability. In order to comprehend the results, the methodology is structured comprehensively. In addition, the use of this methodology ensures a robust and scientific base for results. The adversary models are fuelled by artificial intelligence and machine learning and make cyber-attacks automated, intelligent, and sophisticated. Many modelling techniques adopt complex algorithms and decision-making logics that are difficult to understand in order to monitor and mitigate the immense complexity of emerging cyber threats. Given this backdrop, this research uses agentic AI, with an emphasis on building models with stated and transparent intentions. It was anticipated that the agent-based defense system would increase understanding of how the model of the adaptive adversary operates. The adaptive agent is powered by rules, actions, and intentions previously defined and executed in a particular given sequence. The training process facilitated learning of the rewards and penalties obtained by the agent in response to actions or in a specific situation. Assumptions and constraints are determined by the adaptation of the agent model. These criteria work as a blueprint for the logical machine code that will be generated by the statistical model.

Equ 2: Action Space

$$a_t \in \mathcal{A}$$

- Blocking a particular IP or user.
- Deploying an intrusion detection system (IDS).
- Initiating data encryption.
- Disconnecting compromised devices from the network.

4.1. Data Collection and Preprocessing

An Agentic AI and Reinforcement Learning (RL) based framework is studied. RL can be used to automate determining and executing cyber remediation when real-time observations suggest adversary actions in an environment. Specifically, the design and training is performed on two neural network (NN) based models. The first model identifies user accounts, groups and hosts by predicting window event logs of user group membership event type. Given the model, a set of wCyber Threat Response (CTR) Actions is then determined by examining CTR policies. The second RL agent is trained in network intrusion data, watching for cyber adversary actions that occur in the user account-group-host environment. In simulation, observations are provided to the second RL agent and adaptively select discrete-action decisions in real-time, seeking to alter network data without being caught by the first agent for as long as possible. The study shows the trained RL agent outperforms random action, achieving larger rewards and threat actor's survival times.

The Names of both popular and academic journals are specified. Data in its various forms are essential to clarify the environment, specify any randomness, and inform analytic judgments and decisions. While attention has been justifiably devoted to methodological choices, an opportunity to point to the many related considerations constraining or enhancing those choices is taken. Specifically, this paper is devoted to the detailed presentation of the choices made for data collection and preprocessing. As will be evident is an extension of the concerns of the typical methodological 'coda' after detailed presentation of results.

Four categories of data are used: (1) historical cyber threat data; (2) real-time attack information; (3) environment and adversary modeling; and (4) agent observation/action. Various sources are used to gather individual datasets. First, multiple open source malware repositories are utilized, retrieving malicious files of different varieties. Second, the structure and location of various windows event log storage on the file system is described. Third, a list of operations available in the environment: 'net user', 'query user', 'New-LocalGroup', 'Add-LocalGroupMember', 'query localgroup', and 'New-LocalUser'. Finally, information about time in the simulator Monday is given. To part with, experimental methodology entails preprocessing of such diverse datasets, should a brief summary of each data's statistical properties be found in a tabular format. Data cleansing, data normalization, and time-scale transformation are conducted prior to simulation and modeling, ensuring consistency and quality in the datasets. Since no data preprocessing is identical across different types of data, detailed algorithms for each type are provided. While commonplace missing data and noise handling techniques are used to some extent within preprocessing, typically they are less focused on. For this reason, a 'best practices' approach is taken for each category of datasets to maximize subsequent model accuracy. Methods for dataset storage and retrieval are also a focus since each of the 26 datasets is moderately large.

4.2. Model Development and Training

In subsection, the complexities in conducting model development and training processes are explored along with various considerations that rise. Selection criteria for reinforcement learning algorithms were made meticulously to ensure initial hyperparameter tuning and effective training. Besides, methodologies accounted for in the fine-tuning of model parameters, including in the enhancement of model performance and robustness, are described. Equally, the training strategies applied, incorporating supervised, unsupervised, and semi-supervised learning approaches, are discussed, demonstrating their capacity for adaptation to different scenarios. In addition, the methodologies used to evaluate the performance of trained models are elaborated, ensuring the process is iterative and allowing for different results. Exemplified use cases as to how validated models range up against real-world scenarios are presented, strengthening credibility, as well as issues arising during model development, and so how said issues are tackled.

Model development and, following this, model assessment are possibly those most complicated tasks pertaining to this research undertaking. In accordance with scientific explanation, the collection and sanitisation of feature data accompanying a unified knowledge graph are detailed upfront, prior to even casting a glimpse at machine learning-based processes. The data sets used for model training and accreditation, together with a graphical flow map of the full methodology are displayed. Transparently, these two sections are essential as a grounding bed to ensure understanding regarding the additional methodology used. Conducting model development remains a multifaceted process, and so the computational set-up required is accounted for. Each approach used in developing models to automate cyber threat response entails a sequence of steps, including a selection of uninstructed, semi-supervised, and reinforcement learning strategies. With each strategy, this process concludes with model training, implementation, and post-deployment measures. Due diligence is undertaken to consistently assess model performance, and so these evaluations are detailed within a dedicated sub-section rather than described under individual model development strategies. This treatment necessitates that this sequence will be repetitive, however it is the most logical arrangement. Concluding this section, a brief matter of how this full model development loop is iterative, and iterative improvement is undertaken beyond this methodical write-up is given.

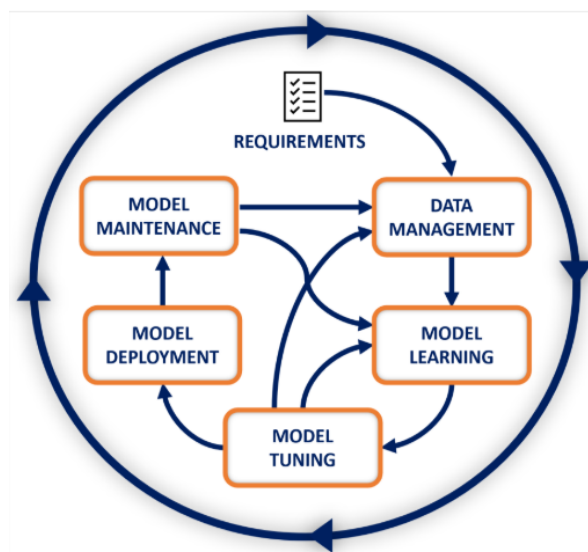


Fig 4: Model Development and Training

5. Case Studies and Applications

This section presents multiple case studies denoting the implementation and possible results and outcomes of various applications of agentic AI and reinforcement learning in cybersecurity circumstances. The case studies here focus exclusively on automated response systems for cyberthreats, as this demonstrates the most advanced, novel, and intricate combination of AI and reinforcement learning capabilities to date in this field.

A case study on the typical application of automated vulnerability reinforcement detection with reinforcement learning and threat response with agentic AI in a moderately large healthcare chain is first examined. Following this, three case studies examining different scenarios for incident response with an increasing level of system independence and complexity, beginning with a small-medium-sized law firm, followed by a large non-profit sector, and concluding with a multinational corporation.

Case Information: Every day companies are being targeted by advanced persistent threats (APTs) and zero-days which are both difficult to detect and respond in a timely manner. Since companies have little to no choice of whether to use information technology for their operations, they are in a vulnerable position and are dependent on their information systems which have an increasing demand to be secure. Companies are thus inclined to

facilitate a safeguard of their data whether personal, customer, business capacities, or intellectual property.



Fig 5: Applications of Automating Cyber Threat

6. Conclusion and Future Directions

In the context of the study, agentic AI systems and reinforcement learning-based technologies are introduced. These technologies contribute to swiftly accommodating the evolving nature of cyber threats and enhancing cybersecurity responses, reducing the discrepancy between the speed of threat response construction and incidents. Agentic AI frameworks take control in cyber threat responses by analyzing diverse possibilities to harness the emerging trends among malware and bots, before instructing their counterparts. Reinforcement learning outcomes are incorporated in the threat response of the cyber defense arrangement through the recommendation of the optimal countermeasure, surpassing manual policy drafting. Additionally, future research directions and open challenges are deliberated as guidance for continuous studies. Due to the multifaceted, thorough, and time-consuming nature of cyber defense, required countermeasures typically form later than threats actually materialize. This discrepancy challenges the timeliness of cyber defense arrangements, leaving them vulnerable to attacks. It is broadly recognized that organizations and societies will not be able to totally eliminate cyber threats. Therefore, improving the resiliency of cyber defense and threat response mechanisms has become the core of the research in preventing or markedly lessening the impacts of cyber-threat incidents. This research aims to construct a cyber defense arrangement that automatizes the response to cyber threats with the contemplated consideration of threatened assets and evolving threats.

With the explosive expansion of internet users and devices, an extensive range of cyber infrastructures has been established to accommodate increasing digital consumption and maintain ubiquitous communication. The soaring dependence on cyber-physical systems and the internet build complex networked systems with elaborate interconnections from day-to-day tasks and industrial processes to governance preparation and military applications.

Equ 3: Q-Learning (Reinforcement Learning)

.

Where:

$$Q(s_t, a_t) = R(s_t, a_t) + \gamma \max_{a_{t+1}} Q(s_{t+1}, a_{t+1})$$

- γ is the discount factor (how much future rewards
- s_{t+1} is the next state after taking action a_t .

6.1. Ethical and Legal Considerations

Cybersecurity is facing a significant challenge that is only expected to grow in scope over the next few years. Agentic artificial intelligence (AI) is expected to address this challenge as it becomes more difficult for a human to specify a reward function for an AI. There are risks and ethical implications of using agentic AI in this context. A wide-ranging regulatory framework exists for AI and autonomous systems in many jurisdictions. Companies and developers are legally responsible for the products when they do not follow regulations or due diligence laws. There are also general problems with interpreting traditional AI and machine learning models. For instance, it is difficult to verify that a particular output is

predictable. Regulation currently complicates the implementation of Deep Reinforcement Learning (DRL) and ARL algorithms. Design or deployment of AIs allowing attackers to interact with the environment in a manner not intended by the designer could be considered harmful by policing or regulating forces. Additionally, bias and other deficiencies in AI prediction models have grown to be an impediment to the deployment of AI for strategic purposes. One final setting that complicates the deployment of automated decision-making systems is the adversarial nature of the environment in which such a system must operate. Cyber threats are continuously evolving and pose a serious danger to individuals and business, and the global economy thus relies on robust cybersecurity defenses. However, it is foreseen that the shortage of cybersecurity talent means that there are not enough experts to defend effectively against the threats. There are also growing markets for tools that make even the unskilled user capable of carrying out cyber-attacks. The number and capabilities of the attacks will increase, and without drastic changes to the field, it will be difficult to counter this.

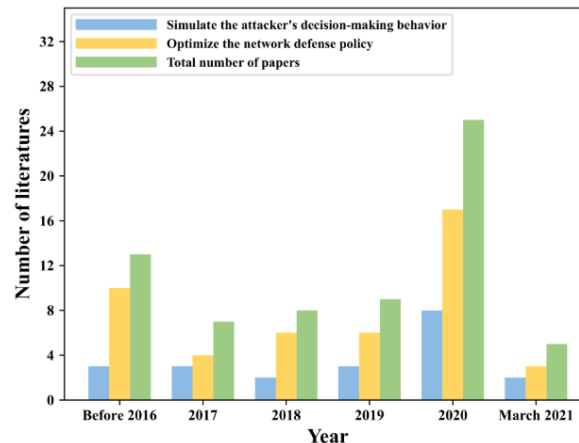


Fig : Reinforcement Learning in Cyber Defense Decision-Making for Intranet Security

6.2. Scalability and Performance Issues

For a long time, the global security community has been promoting detection and collaboration as the primary ways to address escalating cybersecurity threats. However, vigilant observation by security professionals is not feasible with the increasing volume, complexity, and speed of modern IoD/IoT technology. Recognizing this trend, security companies and organizations increasingly embrace the use of artificial intelligence (AI), machine learning, and reinforcement learning (RL) techniques. These agentic AI-based technologies allow an entity to detect and respond to threats faster. However, the wider implementation of agentic AI in cybersecurity contexts has encountered hurdles, particularly when applied to small and medium-sized enterprises with limited IT infrastructure and resources.

Regarding scalability, the structure, application, and size of organizations' computing infrastructure can significantly affect an agentic AI system's utility and effectiveness. In terms of RL, the rapid complex decision-making regarding network firewall rules requires more extensive computation, while available resources and responsiveness are challenged to keep up. Smaller organizations with limited resources can only manage slow computation. When the network infrastructure grows rapidly, processing speed needs to be optimized to manage resources effectively. However, speeding up the resource-hungry search processes for computing optimal IP allow/block rules will be a non-trivial task. One possible approach is to move the computation to the cloud, and another is to adopt a modular architecture to plunge the most computationally intensive tasks into isolated computational parts. Both can be chosen to work in a complementary way once the scale of the infrastructure grows to a level when it becomes more effective. Local machine-based DPI is recognized as inefficient and is particularly taxing on resources. Moreover, neither the network nor security professionals can appropriately be explained so that no insights are directly applicable.

7. References

- [1] Sikha, V. K. Ease of Building Omni-Channel Customer Care Services with Cloud-Based Telephony Services & AI.

- [2] Nampally, R. C. R. (2021). Leveraging AI in Urban Traffic Management: Addressing Congestion and Traffic Flow with Intelligent Systems. In *Journal of Artificial Intelligence and Big Data* (Vol. 1, Issue 1, pp. 86–99). Science Publications (SCI-PUB). <https://doi.org/10.31586/jaibd.2021.1151>
- [3] Ganesan, P., Sikha, V. K., & Siramgari, D. R. TRANSFORMING HUMAN SERVICES: LEVERAGING AI TO ADDRESS WORKFORCE CHALLENGES AND ENHANCE SERVICE DELIVERY.
- [4] Vankayalapati, R. K., & Syed, S. (2020). Green Cloud Computing: Strategies for Building Sustainable Data Center Ecosystems. *Online Journal of Engineering Sciences*, 1(1), 1229. Retrieved from <https://www.scipublications.com/journal/index.php/ojes/article/view/1229>
- [5] Eswar Prasad Galla.et.al. (2021). Big Data And AI Innovations In Biometric Authentication For Secure Digital Transactions Educational Administration: Theory and Practice, 27(4), 1228 –1236 Doi: 10.53555/kuey.v27i4.7592
- [6] Sikha, V. K. INTELLIGENT SYSTEMS AND APPLICATIONS IN ENGINEERING.
- [7] Ganesan, P., & Sanodia, G. (2023). Smart Infrastructure Management: Integrating AI with DevOps for Cloud-Native Applications. *Journal of Artificial Intelligence & Cloud Computing*. SRC/JAICC-E163. DOI: doi.org/10.47363/JAICC/2023 (2) E163 J Arti Inte & Cloud Comp, 2(1), 2-4.
- [8] Mohit Surender Reddy, Manikanth Sarisa, Siddharth Konkimalla, Sanjay Ramdas Bauskar, Hemanth Kumar Gollangi, Eswar Prasad Galla, Shravan Kumar Rajaram, 2021. "Predicting Tomorrow's Ailments: How AI/ML Is Transforming Disease Forecasting", *ESP Journal of Engineering & Technology Advancements*, 1(2): 188-200.
- [9] Janardhana Rao Sunkara, Sanjay Ramdas Bauskar, Chandrakanth Rao Madhavaram, Eswar Prasad Galla, Hemanth Kumar Gollangi, Data-Driven Management: The Impact of Visualization Tools on Business Performance, *International Journal of Management (IJM)*, 12(3), 2021, pp. 1290-1298. <https://iaeme.com/Home/issue/IJM?Volume=12&Issue=3>
- [10] Ganesan, P. (2021). Advanced Cloud Computing for Healthcare: Security Challenges and Solutions in Digital Transformation. *International Journal of Science and Research (IJSR)*, 10(6), 1865-1872.
- [11] Vaka, D. K. (2020). Navigating Uncertainty: The Power of 'Just in Time SAP for Supply Chain Dynamics. *Journal of Technological Innovations*, 1(2).
- [12] Gagan Kumar Patra, Chandrababu Kuraku, Siddharth Konkimalla, Venkata Nagesh Boddapati, Manikanth Sarisa, An Analysis and Prediction of Health Insurance Costs Using Machine Learning-Based Regressor Techniques, *International Journal of Computer Engineering and Technology (IJCET)* 12(3), 2021, pp. 102-113. <https://iaeme.com/Home/issue/IJCET?Volume=12&Issue=3>
- [13] Ganesan, P. (2021). Leveraging NLP and AI for Advanced Chatbot Automation in Mobile and Web Applications. *European Journal of Advances in Engineering and Technology*, 8(3), 80-83.
- [14] Venkata Nagesh Boddapati, Eswar Prasad Galla, Janardhana Rao Sunkara, Sanjay Ramdas Bauskar, Gagan Kumar Patra, Chandrababu Kuraku, Chandrakanth Rao Madhavaram, 2021. "Harnessing the Power of Big Data: The Evolution of AI and Machine Learning in Modern Times", *ESP Journal of Engineering & Technology Advancements*, 1(2): 134-146.
- [15] Ganesan, P. (2021). Cloud Migration Techniques for Enhancing Critical Public Services: Mobile Cloud-Based Big Healthcare Data Processing in Smart Cities. *Journal of Scientific and Engineering Research*, 8(8), 236-244.
- [16] Kothapalli Sondinti, L. R., & Syed, S. (2021). The Impact of Instant Credit Card Issuance and Personalized Financial Solutions on Enhancing Customer Experience in the Digital Banking Era. *Universal Journal of Finance and Economics*, 1(1), 1223. Retrieved from <https://www.scipublications.com/journal/index.php/ujfe/article/view/1223>
- [18] Mandala, V., & Surabhi, S. N. R. D. Intelligent Systems for Vehicle Reliability and Safety: Exploring AI in Predictive Failure Analysis.

- [19] Vankayalapati, R. K., & Rao Nampalli, R. C. (2019). Explainable Analytics in Multi-Cloud Environments: A Framework for Transparent Decision-Making. *Journal of Artificial Intelligence and Big Data*, 1(1), 1228. Retrieved from <https://www.scipublications.com/journal/index.php/jaibd/article/view/1228>
- [20] Ganesan, P. (2020). Balancing Ethics in AI: Overcoming Bias, Enhancing Transparency, and Ensuring Accountability. *North American Journal of Engineering Research*, 1(1).