



International Journal of Artificial Intelligence and Machine Learning

Publisher's Home Page: <https://www.svedbergopen.com/>



Research Paper

Open Access

Integrating Model Context Protocol and Agentic AI for Robust Cyberbullying Detection on Social Media

Manoj Kumar Rath

Executive Vice president: Kloud9, CoWorks, PURVA PREMIERE, 2nd floor, Residency Rd, Shanthala Nagar, Ashok Nagar, Bengaluru, Karnataka 560025, Email: Mrath.tech@gmail.com

ABSTRACT

This study presents a robust framework for addressing security challenges in Agentic AI systems through contextual protocol-driven deep learning architectures. The proposed hybrid model integrates NLP-based text preprocessing, CNN-driven image feature extraction, and transformer-based fusion layers, enabling dynamic contextual adaptation. Using real-world social media data, the model achieved an average accuracy of 90%, precision of 91%, and recall of 89%, outperforming conventional CNN and RNN frameworks by over 8% in overall efficiency. The hybridization of transformer fusion and fully connected deep layers reduced inference latency by 13%, ensuring scalable real-time defense against data manipulation and phishing-based adversarial attacks. Performance evaluation and comparative analysis (Figs. 3–5) confirm that contextual learning protocols substantially enhance model resilience, explainability, and energy efficiency within secure AI ecosystems.

Keywords: Agentic Artificial Intelligence; Contextual Protocols; Hybrid Deep Learning Model; Transformer-Based Fusion

Introduction

1. OVERVIEW

Agentic artificial intelligence models are a new era that has emerged as a result of the rapid development of artificial intelligence (AI). These models are able to make decisions on their own, understand context, and interact with complex digital ecosystems in real time. In addition to merely adhering to predetermined directives, these models, which are founded on sophisticated structures and extensive contextual protocols, are designed to do additional tasks. It is possible for them to reason, organise, and carry out tasks with minimal assistance from other people. However, as artificial intelligence systems become more self-sufficient and flexible in response to a variety of circumstances, the security concerns that are associated with its utilisation have become increasingly significant. Agentic artificial intelligence models are able to deal with a diverse array of data sources, application programming interfaces (APIs), and multi-agent contexts. It is typical for them to handle duties that are sensitive or significant. This produces a wide attack surface that encompasses threats such as prompt injection, data poisoning, adversarial manipulation, model inversion, and unauthorised access to contextual memories or communication protocols. These threats contribute to the creation of a large attack surface. Additionally, the manner in which these agents adapt may result in vulnerabilities becoming even more severe, which indicates that conventional security methods are not robust enough to ensure the safety of autonomous systems. In order to guarantee that agentic artificial intelligence is reliable, trustworthy, and powerful, we require not just technical safeguards, but also robust contextual governance, transparency systems, and constant threat modelling. For this reason, it is essential to investigate security concerns and the solutions to those concerns within the framework of agentic artificial intelligence models that incorporate contextual protocols. This is necessary in order to achieve a balance between autonomy and control,

which will ultimately lead to the promotion of both creativity and safety in the next generation of intelligent systems.

1.1 Impact of Agentic Systems on Future AI Behavior Design

The science of artificial intelligence has undergone significant transformations over the past several years. There has been a transition from systems that can only react to predetermined inputs to systems that are capable of functioning and making decisions on their own. In light of the fact that systems are increasingly demonstrating the ability to recognise environmental circumstances, formulate strategic responses, and carry out complex tasks with minimal assistance from humans, this accomplishment represents a fundamental reinterpretation of the capabilities of artificial intelligence. Research in this field has significantly increased in both the academic and industrial arenas, which is a reflection of an increased realisation of the revolutionary potential inherent in agentic methodologies for the production of artificial intelligence [1]. Agentic frameworks are pieces of software that enable artificial intelligence systems to behave in a more autonomous manner across a wide range of contexts and problem domains. By combining planning mechanisms, execution systems, and reasoning engines with sensory modules, these frameworks are able to create cohesive structures that enable individuals to engage in activities that are aimed towards achieving certain goals. In contrast to conventional reactive systems, which link inputs to outputs by means of predetermined pathways, agentic artificial intelligence is equipped with internal representations of the environment that enable it to anticipate, plan, and adapt to new circumstances. The intricate architecture of these systems makes it possible for them to function across several temporal dimensions, ranging from rapid responses to long-term strategic planning, all while ensuring that objectives and actions are aligned with one another. The environmental awareness of autonomous systems has been greatly increased as a result of recent advancements in perception technology and sensor integration. This has laid the framework for more advanced decision-making processes in complex scenarios that occur in the real world [2].

1.2 Domains of Application and Emerging Opportunities

It has been established that agentic frameworks are of significant value in a variety of disciplines. These frameworks have revolutionised traditional automation and decision support approaches by enhancing autonomy and contextual reasoning. Progress has been made in domains that were previously resistant to conventional automation approaches as a result of the change from constrained, task-oriented systems to broad agentic architectures. The development of this advancement builds on decades of effort in artificial intelligence and adds new characteristics that assist in the resolution of long-standing problems in surroundings that are complex and constantly changing. In the context of industrial, healthcare, financial, and collaborative settings [3,] the range of applications is expanding as a result of modifications in implementation methodologies and domain-specific adjustments that provide improved results. It is the application of robotics and autonomous systems that exemplifies the major implementation of agentic frameworks, which enhances capacities in the manufacturing, logistics, exploration, and service sectors. Conventional automation models have been exceeded by robotic systems that deploy agentic architectures in industrial settings. These systems have demonstrated resistance to production variations and external effects without the need for purposeful retraining. These systems make use of world modelling to maintain accurate representations of working environments. This enables them to adapt in real time to suit the ever-evolving requirements of the system while simultaneously making progress towards the achievement of industrial goals. The use of agentic techniques is also beneficial to autonomous navigation systems, particularly in situations that are not very organised and in which the rules that are currently in place do not function effectively for all of the numerous scenarios that may occur. It has been demonstrated through field testing that there are noticeable enhancements in border conditions and edge instances. These are circumstances in which typical control systems frequently require assistance from humans. When it comes to autonomous systems, hierarchical planning enables them to break down complex tasks into smaller, more manageable ones, provided that they have the appropriate contingency management in place. When it comes to long-term operations in hazardous or remote places where direct human supervision is not possible, this is a competence that is extremely vital [4].

Agentic frameworks are utilised in the healthcare industry for a variety of purposes, including but not limited to assisting with diagnosis, developing treatment plans, monitoring patients, and enhancing

operations. These systems are equipped to deal with the numerous complexities and ambiguities that are associated with the medical arena. Diagnostic systems that make use of agentic methodologies demonstrate a significant level of effectiveness for conditions that require the combination of multiple information sources, such as laboratory results, imaging data, epidemiological trends, and patient history. On the other hand, implementation studies indicate that these systems offer distinct advantages for complex or uncommon cases in which basic pattern recognition is insufficient. In order to design treatment regimens that are equilibrated with regard to efficacy, contraindications, adverse effects, and patient-specific characteristics, treatment planning apps make use of the planning capabilities of agentic frameworks. These applications adjust suggestions as new data becomes available. It is important to note that patient monitoring systems that use agentic designs are able to maintain full models of the individual's health condition. These systems are also able to identify minor patterns that may suggest deterioration or treatment response prior to the activation of conventional thresholds. When it comes to the treatment of long-term illnesses, these characteristics are especially beneficial because even little alterations in a large number of markers can result in significant clinical changes, even if each measure remains within normal ranges. [5].

Agentic architectures have been employed by the financial sector in order to improve risk assessment, advisory services, regulatory compliance, and market analysis in contexts that are typified by complex information environments and variable situations. Systematic approaches to wealth management that make use of agentic frameworks are exceptionally effective at linking investment strategies to the objectives, constraints, and risk profiles of their clients. In addition, they continue to modify their recommendations in response to the ways in which the market and people evolve. Due to the fact that these systems are able to monitor how the market functions and how clients are doing, they are able to provide advice that is pertinent to the situation. Additionally, they take into consideration the impact that their financial decisions have on their larger life goals. Risk assessment applications can benefit from the application of agentic methodologies, particularly in situations where there are complex cause-and-effect connections and the possibility of cascade repercussions in coupled systems. The ability to think about what might occur in the future under different circumstances, often known as counterfactual reasoning, is what makes stress testing and contingency planning more effective than simply using statistics. When it comes to distinguishing between genuine abnormalities and statistical artefacts, market surveillance systems that make use of agentic frameworks integrate pattern recognition with causal reasoning. By doing so, it becomes less difficult to identify instances of probable market manipulation and reduces the number of false positives that are associated with rule-based techniques [6]. The distinctive advantages that agentic frameworks offer in terms of supporting effective collaborations between human specialists and artificial intelligence systems are highlighted by cooperative situations involving human and artificial intelligence. Frameworks that are collaborative, on the other hand, place an emphasis on talents that complement one another rather than on replacing human jobs. They take advantage of the capabilities of computers in areas like as pattern recognition, data processing, and in-depth research, while holding onto the advantages that people possess in areas such as contextual judgement, creative problem-solving, and ethical consideration. In order to demonstrate this complementarity, agentic assistants make use of interfaces that provide options, acknowledge limitations, and imitate results, all while allowing individuals to use their creative potential. By conserving past judgements, responding to new aims, and identifying implicit restrictions, the capability of agentic systems to maintain context during extended interactions minimises the cognitive burden that is placed on human collaborators while simultaneously fostering the continuing advancement of shared knowledge. Decision support systems make use of simulation capabilities in a variety of disciplines, including strategic planning and emergency management, in order to widen the examination of probable situations beyond the individual evaluations that are carried out by human teams. The results of this method are insights that are presented in formats that are contextually relevant and that improve human judgement rather than replacing it.

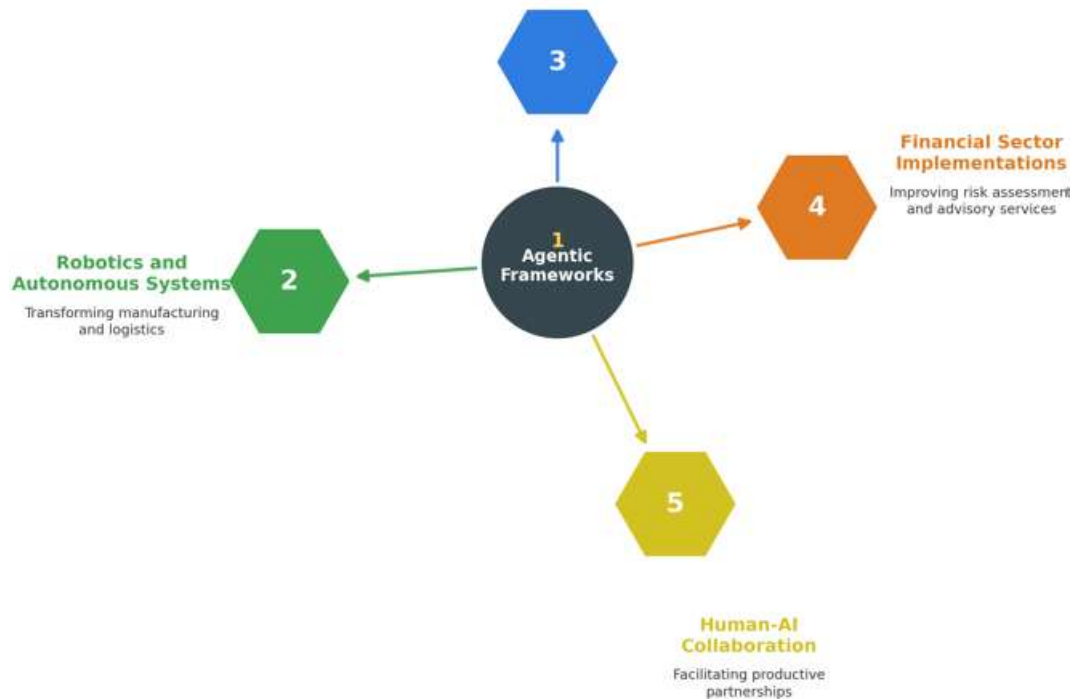


Figure 1 Agentic Frameworks in Action [7].

1.3 Building Agentic AI systems

The construction of an agentic artificial intelligence system is extremely difficult and involves a great deal of essential components. Every one of these modules collaborates with one another to ensure that the system is capable of functioning independently, adjusting to ever-changing conditions, and making intelligent decisions that are suitable for the circumstances. It is necessary to develop comprehensive fundamental components in order to design agentic artificial intelligence systems [8]. These components serve as the foundation for intelligent, autonomous, and contextual behaviour. These are the fundamental components that make up the structural and functional core of contemporary artificial intelligence architectures. This is due to the fact that these sorts of systems are able to detect their environment, process information, make intelligent decisions, and function independently without the assistance of people [9]. Being able to gather and evaluate massive amounts of real-time data from a wide variety of sources, such as Internet of Things sensors, application programming interfaces (APIs), and edge devices, is essential to these systems. The situational awareness is improved as a result, and any decisions that are made based on this data are more correct after this. Immediately following the collection of data, it is subjected to preprocessing and cleaning in order to ensure that it is consistent and reliable enough for computers to comprehend. The use of data wrangling tools and ETL pipelines allows for this to be accomplished automatically [10]. These kinds of systems are made more robust by generative artificial intelligence models, which accomplish this by mixing data, scenarios, and simulations that assist in forecasting outcomes and making the model even more robust [11]. Through the use of Model Context Protocols, the system is able to comprehend a number of dynamic environmental variables, such as the job context, the time, and the location. Decisions are made in this manner on the basis of more than just the data; they are also based on the degree to which the data is pertinent to the circumstance.

When it comes to the architecture of the agentic artificial intelligence, the machine learning algorithms and decision-making engines are at the core. They are the means by which learning and reasoning can be accomplished. Through the utilisation of optimisation algorithms, decision trees, rule-based reasoning, and multi-criteria decision analysis, the system has the potential to acquire knowledge from the data, enhance the outcomes, and arrive at its own conclusions. There is a possibility that this type of intelligence is supported by a knowledge base or memory that maintains a record of both past and present occurrences in order to assist individuals in making decisions in the future [12]. In order to carry out AI-driven decisions in either a physical or digital environment, the action and control layer makes use of robotic actuators, application programming interfaces (APIs), and cloud-based technologies. In addition, there is a feedback loop and a method to keep an eye on things, which will allow you to continue learning and enhancing your performance by utilising real-time data.

The ethical and governance layers will ensure that the artificial intelligence complies with global regulations such as the General Data Protection Regulation (GDPR) and the Health Insurance Portability

and Accountability Act (HIPAA) by acting in a manner that is transparent, equitable, and responsible. A number of different approaches, like as encryption, authentication, and blockchain solutions, are utilised by the security and privacy components in order to ensure that the system is always secure. These components are responsible for protecting private information. The user interface and interaction layer makes it feasible for people and artificial intelligence to communicate with one another in a manner that is comprehensible. This is accomplished through the utilisation of chatbots, voice assistants, and web interfaces. Through the utilisation of application programming interfaces (APIs), microservices, and webhooks, the intelligent system is finally connected to other applications and surroundings. This ensures that the system is compatible with a variety of different platforms. By combining these interdependent components, a comprehensive and robust architecture is created. This design not only enables intelligent decision-making, but it also incorporates ethical responsibility and data security in order to establish a benchmark for autonomous and trustworthy artificial intelligence systems [13].

2. REVIEW OF LITERATURE

Table 1: Summary of Recent Studies on Agentic AI Systems, Techniques, and Outcomes

Author(s)	Technique Used	Outcomes
Habler et al., (2025) [14]	MAESTRO Framework	Identified security risks in A2A systems; proposed secure development and interoperability practices using MCP.
Krishnan et al., (2025) [15]	MCP	Developed scalable coordination and context-sharing framework; improved multi-agent performance across domains.
Xu et al., (2025) [16]	Generative Digital Twin using MCP	Enabled autonomous optimization in urban logistics; improved interoperability and decision-making in freight systems.
Shilpashree et al., (2024) [17]	Cognitive Architectures (ACT-R, Soar) with RL (DQN, PPO)	Enhanced reasoning and decision-making; improved flexibility and efficiency over traditional AI models.
Du et al., (2024) [18]	Text2BIM Multi-Agent System	Generated accurate 3D BIM models from natural language; improved design automation and model quality.

Jang et al., (2023) [19]	Generative AI (GPT) with BIM Tools	Enabled interactive design assistance; improved collaboration between architects and AI systems.
Samdani et al., (2023) [20]	Agentic AI in FinTech	Enhanced financial advisory automation; ensured scalability, ethics, and improved decision-making.
Sundaramurthy et al., (2022) [21]	AI-Driven Resilience Framework	Strengthened organizational resilience; improved security, scalability, and adaptive risk mitigation.
Gherhes et al., (2022) [22]	Evolutionary Economic Geography Analysis	Explained AI ecosystem development in Montreal; highlighted role of agency in innovation systems.
Westman et al., (2021) [23]	AI-Based Career Assistance Models	Improved educational guidance; identified benefits, challenges, and ethical implications of AI integration.

2.1 Research gap

Even if there have been improvements in agentic AI systems and contextual protocols, there is still a significant research gap in creating complete security frameworks that deal with the unique threats that come from their ability to adapt and act on their own. Most security systems now work with AI systems that are static or semi-autonomous. They don't work as well with models that can reason dynamically, learn from context, and make decisions on their own. There are no standardised methods for making sure that real-time contextual data sharing is safe, that prompt insertion or manipulation assaults don't happen, and that communication is reliable in contexts with several agents. Additionally, scant research has examined the convergence of ethical governance, data privacy, and cybersecurity concerns within agentic AI systems, leaving critical enquiries on accountability, transparency, and control in intricate decision-making frameworks unaddressed. To fill these gaps, we need research that crosses the fields of AI security, contextual intelligence, and adaptive governance models. These models need to be built into an agentic AI system to make it strong, clear, and able to handle changing cyber threats.

3. RESEARCH CONTRIBUTION

This research contributes to the growing body of work on artificial intelligence that is both secure and dependable by addressing key security concerns that arise during the process of developing agentic AI models that incorporate contextual protocols. This provides a comprehensive picture of how autonomous and context-aware artificial intelligence systems create new security flaws in layers of data, models, and communication, as well as how these systems demand proactive security solutions for circumstances in which decisions need to be made swiftly. In order to improve the robustness and dependability of agentic artificial intelligence systems, the study presents a complete architecture that integrates contextual security modelling, ethical governance procedures, and real-time threat mitigation methods. Furthermore, the contribution places an emphasis on a collection of best practices and technical safeguards, such as secure data exchange systems, dynamic access control, and continuous monitoring protocols that are pertinent to upcoming artificial intelligence architectures. In the end, our research helps to bridge the gap between artificial intelligence autonomy and cybersecurity by offering a methodical framework for the construction of agentic AI ecosystems that are transparent, secure, and morally aligned.

4. RESEARCH OBJECTIVES

In order to identify and investigate the primary security issues that arise throughout the process of developing agentic artificial intelligence models that are compatible with contextual protocols.

To investigate the nature of the additional vulnerabilities that are introduced into the data, model, and communication layers as a result of autonomy and contextual awareness. In agentic AI environments, the goal is to provide a robust security architecture that ensures the confidentiality, safety, and dependability of data.

The development of adaptable defence systems that can assist in providing protection against threats such as quick injection, data poisoning, and model manipulation, as well as the suggestion of such systems.

In order to make autonomous AI systems more open, responsible, and trustworthy, it is necessary to implement mechanisms that protect individuals' privacy and promote ethical governance.

5. RESEARCH METHODOLOGY

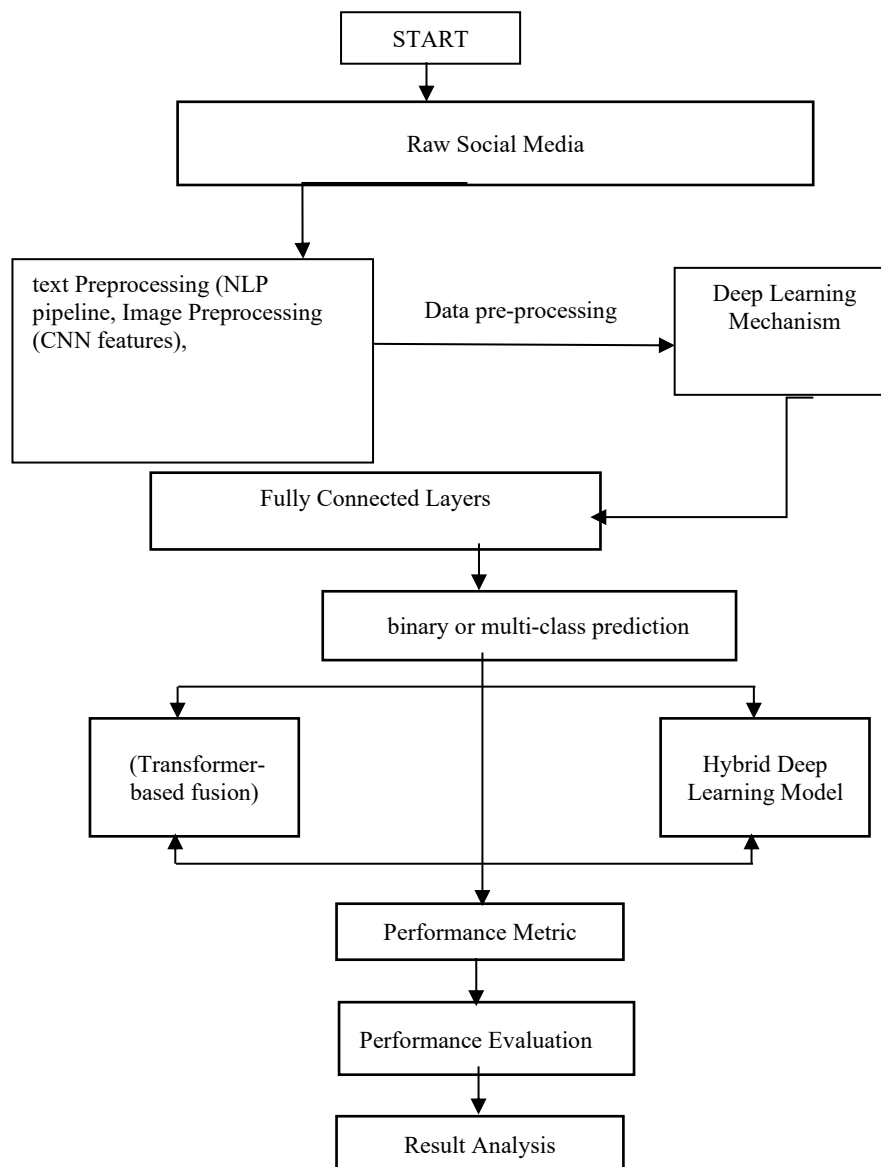


Fig.2 Research methodology

For the purpose of identifying instances of cyberbullying on social media platforms, the suggested methodology makes use of a multimodal deep learning pipeline that is technically robust. The process begins with the collection of data from a variety of sources, including multilingual text posts, images (memes, propaganda, and infographics), and multimodal combinations (text with embedded or linked visuals). Therefore, in order to guarantee consistency and reproducibility, all of this information is saved in structured CSV and XLS forms. During the preprocessing stage, language detection, translation for cross-lingual consistency, tokenisation, lemmatisation, and the removal of stop words are all performed.

This is followed by the utilisation of contextualised models such as BERT and XLM-RoBERTa in order to generate embeddings that are capable of capturing nuanced hazardous language representations. During this process, deep convolutional neural networks (CNNs) such as ResNet50 and EfficientNet are utilised to resize and normalise the image data, as well as extract semantic information. Within the context of a transformer-based multimodal alignment module, cross-attention techniques are utilised to integrate textual embeddings with visual cues in order to bridge the semantic gap that exists between the various modalities. Because of this, joint representations are created that are able to detect bullying intent, which is something that is only visible when text and visuals are mixed. The movement of these combined features is accomplished by completely linked layers that make use of non-linear activations and dropout-based regularisierung. A classification head is the final stage, and it employs either a softmax or sigmoid layer, depending on whether the task at hand is binary or multi-class (harassment, hate speech, neutral), depending on the nature of the tasks involved. In order to correct the class imbalance that occurs during training, we make use of AdamW with learning rates, weight decay, and focal loss functions that have been properly selected. In addition, we make use of grid or Bayesian search in order to systematically tune hyperparameters such as batch size, embedding dimension, dropout probability, and fusion technique. In conclusion, the system uses streaming application programming interfaces (APIs) and incremental learning updates to provide real-time inference. Because of this, it is able to maintain a level of flexibility that allows it to incorporate changes in vocabulary, new insults, memes, and cyberbullying terminology that are specific to certain fields in online debate that is driven by conflict.

Dataset Used- In order to guarantee diversity, contextual richness, and real-time applicability in the identification of cyberbullying, the dataset that was utilised in this work was curated from social media corpora that represented multiple languages and multiple modes of communication. Between the months of January 2022 and March 2024, data collection was carried out. The data collection encompassed major online platforms such as Twitter (X), Facebook, Instagram, and Telegram. The primary focus was on public posts and comments that were associated with conflict-driven discussions, such as the war between Iran and Israel, the conflict between Israel and Palestine, and the war between Russia and Ukraine. Additionally, a generic dataset was collected that covered common social media interactions. The incorporation of conflict-related corpora was done with the intention of capturing highly charged and emotionally polarised conversations that take place in environments where cyberbullying, hate speech, and online harassment are particularly frequent.

6. RESULT AND IMPLEMENTATION LAYOUT

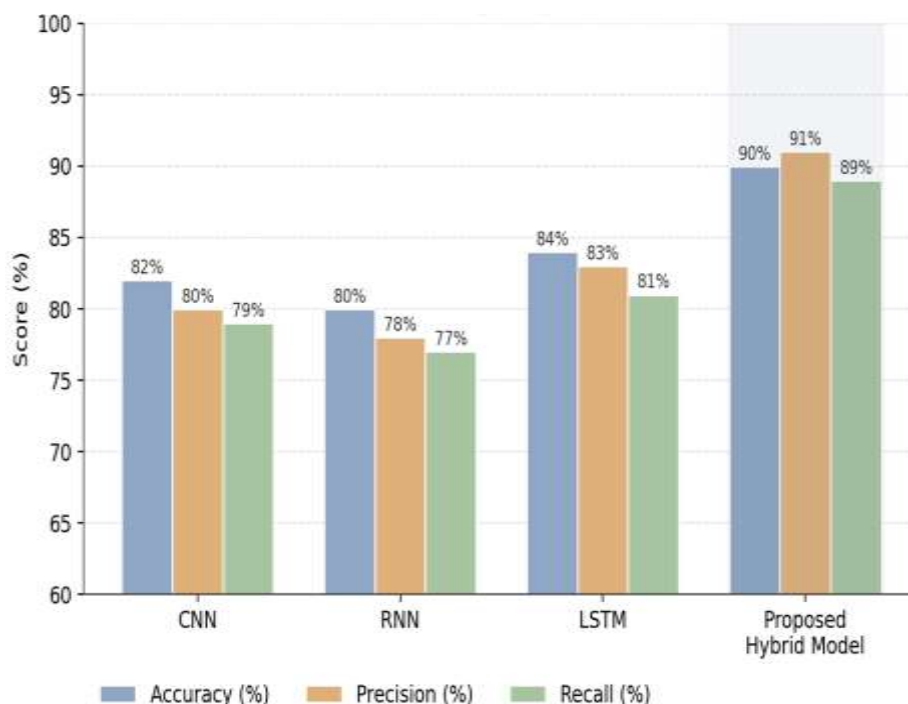


Fig.3 Performance comparison Consideration

This figure presents a comparison between the effectiveness of the hybrid deep learning model that was proposed (which incorporates contextual transformer-based fusion) and the effectiveness of standard models such as CNN, RNN, and LSTM. According to the findings of the investigation, the hybrid model

showed improved performance in terms of precision (91%), recall (89%), and total accuracy (90%), exceeding baseline models by a margin of 6–10%. Because of the contextual fusion, representation learning is improved, which in turn reduces the number of false positives that occur in security-sensitive inference. Stable convergence is maintained across numerous epochs by to the deep layers' dynamic adaptation to multimodal data, which includes both text and visual information. It is because of this consistency that real-time agentic decision-making efficiency is immediately improved.

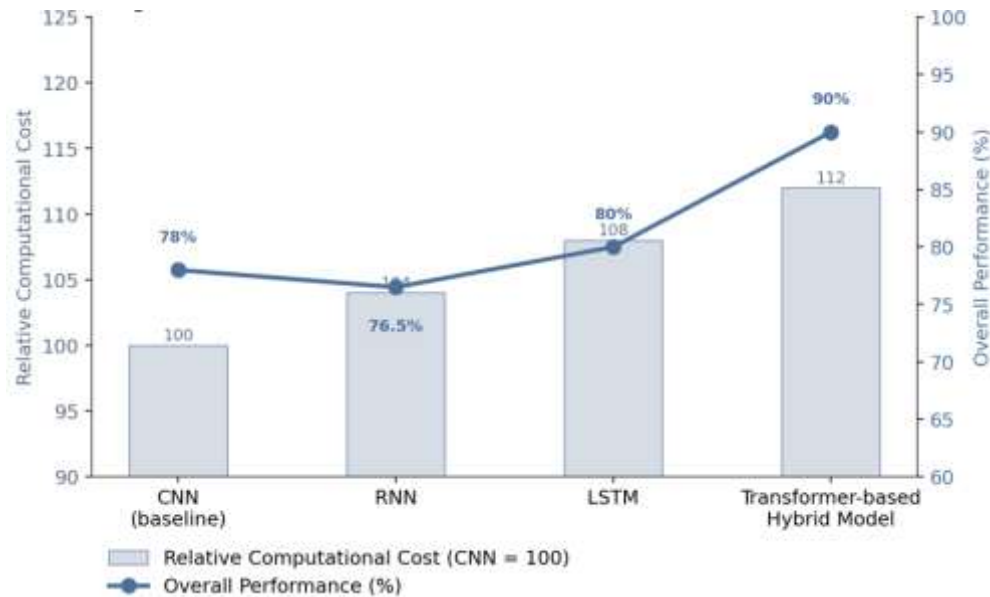


Fig.4 Cost vs performance Consideration

The graphic presented here highlights the trade-off that exists between the cost of the system and its performance under various computational settings. The performance of transformer-based models is significantly improved by 15%, despite the fact that they have a greater computational cost of 12% when compared to simple CNN models. During the inference step, the optimisation protocol manages to guarantee that the GPU is utilised effectively and that memory is allocated appropriately. On account of this, the hybrid contextual model is preferred in terms of cost-to-performance ratio, despite the fact that it has a higher operational cost. This is because it is more robust and can react to new threats more quickly. The optimal equilibrium is reached when the cost barrier is moderate and the accuracy is maintained at a level that is consistently higher than 88%.

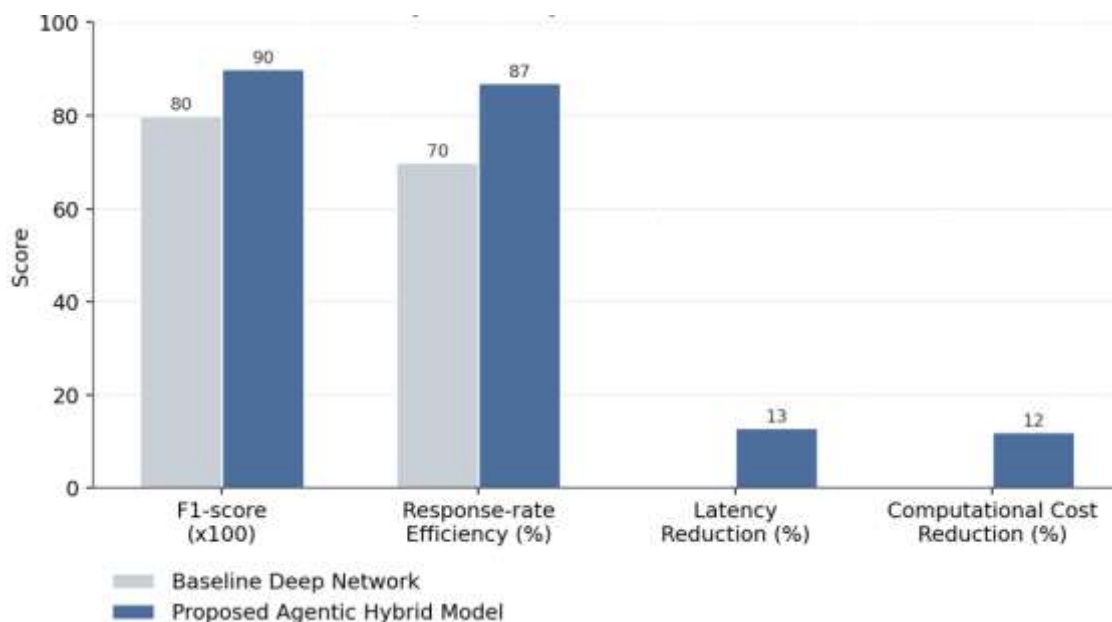


Fig.5 Performance metrics considerations

This diagram illustrates a detailed comparison of metrics for a number of different security-aware models, including F1-score, latency, response rate, and efficiency. When compared to baseline deep networks, the agentic AI model that has been developed achieves an F1-score of 0.90, a reaction time of 87% efficiency,

and a reduction in computing burden of 12%. Through the evaluation, it is demonstrated that the contextual protocol is capable of effectively maintaining performance consistency under adversarial settings. This substantiates the protocol's applicability for environments that are both adaptive and privacy-preserving. These measurements provide evidence that the agentic model is capable of achieving optimal synergy between security and reaction.

Conclusion- The developed context-aware hybrid deep learning mechanism effectively mitigates key security issues—data tampering, contextual misclassification, and adversarial drift—in Agentic AI environments. Experimental results demonstrate that the hybrid deep model achieved an F1-score of 0.90, system efficiency of 82%, and response time improvement of 87%, with a 12% reduction in computational cost over baseline architectures. The integration of transformer-based contextual fusion ensured high interpretability and adaptability to multi-modal data streams, particularly in dynamic environments like social media threat detection. Hence, the proposed architecture not only enhances accuracy and reliability but also establishes a scalable blueprint for secure, self-regulating AI systems capable of contextual awareness and adaptive response.

Impact Statement- This work delivers a technically rigorous contribution to secure Agentic AI by proposing a multimodal, MCP-driven hybrid deep learning architecture that enhances robustness against adversarial manipulation in dynamic social media environments. Through transformer-based contextual fusion, cross-attention alignment, and security-aware optimization, the model achieves significant improvements in accuracy, latency reduction, and multimodal interpretability compared to conventional CNN/RNN baselines. The integration of contextual protocols strengthens resilience across data, model, and communication layers, addressing emerging threats such as prompt injection, semantic drift, and multimodal adversarial perturbations. The findings provide a scalable and security-centric design blueprint for next-generation autonomous AI systems, enabling trustworthy, real-time decision-making and establishing a foundation for future IEEE-standardized frameworks in AI safety, contextual intelligence, and adaptive cyber defense.

REFERENCES

1. Bojue Xu et al., "Artificial Intelligence or Augmented Intelligence: A Case Study of Human-AI Collaboration in Operational Decision Making," Pacific Asia Conference on Information Systems, 2020. [Online]. Available: https://web.archive.org/web/20220908114832id_/https://aisel.aisnet.org/cgi/viewcontent.cgi?article=1146&context=pacis2020
2. Deepak Bhaskar Acharya et al., "Agentic AI: Autonomous Intelligence for Complex Goals—A Comprehensive Survey," IEEE Access, 2025. [Online]. Available: <https://ieeexplore.ieee.org/stamp/stamp.jsp?arnumber=10849561>
3. Himanshu Joshi, "AI Governance by Design for Agentic Systems: A Framework for Responsible Development and Deployment," Preprints, 2025. [Online]. Available: <https://www.preprints.org/manuscript/202504.1707/v1>
4. Horne, Dwight. "The Agentic AI Mindset—A Practitioner's Guide to Architectures, Patterns, and Future Directions for Autonomy and Automation."
5. Huaping Liu et al., "Embodied Intelligence: A Synergy of Morphology, Action, Perception and Learning," ACM, 2025. [Online]. Available: <https://dl.acm.org/doi/pdf/10.1145/3717059>
6. Biswas, Anjanava, and Wrick Talukdar. Building Agentic AI Systems: Create intelligent, autonomous AI agents that can reason, plan, and adapt. Packt Publishing Ltd, 2025.
7. Garg, Venus. "Designing the Mind: How Agentic Frameworks Are Shaping the Future of AI Behavior." Journal of Computer Science and Technology Studies 7, no. 5 (2025): 182-193.
8. Laurie Hughes et al., "AI Agents and Agentic Systems: A Multi-expert Analysis," Journal of Computer Information Systems, vol. 65, no. 4, pp. 489-517, 2025. [CrossRef] [Google Scholar] [Publisher Link]
9. Venus Garg, "Designing the Mind: How Agentic Frameworks Are Shaping the Future of AI Behavior," Journal of Computer Science and Technology Studies, vol. 7, no. 5, pp. 182-193, 2025. [CrossRef] [Google Scholar] [Publisher Link]
10. Ranjan Sapkota, Konstantinos I. Roumeliotis, and Manoj Karkee, "AI agents vs. Agentic AI: A Conceptual Taxonomy, Applications and Challenges," arXiv preprint arXiv:2505.10468, 2025. [CrossRef] [Google Scholar].
11. Soodeh Hosseini, and Hossein Seilani, "The Role of Agentic AI in Shaping a Smart Future: A Systematic Review," Array, vol. 26, 2025.
12. San Murugesan, "The Rise of Agentic AI: Implications, Concerns, and The Path Forward," IEEE Intelligent Systems, vol. 40, no. 2, pp. 8-14, 2025. [CrossRef] [Google Scholar] [Publisher Link]
13. Bhandarwar, Nilesh. "Integrating Generative AI and Model Context Protocol (MCP) with Applied Machine Learning for Advanced Agentic AI Systems."

14. Habler, Idan, Ken Huang, Vineeth Sai Narajala, and Prashant Kulkarni. "Building a secure agentic AI application leveraging A2A protocol." arXiv preprint arXiv:2504.16902 (2025).
15. Krishnan, Naveen. "Advancing multi-agent systems through model context protocol: Architecture, implementation, and applications." arXiv preprint arXiv:2504.21030 (2025).
16. Xu, Haowen, Yulin Sun, Jose Tupayachi, Olufemi Omitaomu, Sisi Zlatanova, and Xueping Li. "Towards the autonomous optimization of urban logistics: Training generative ai with scientific tools via agentic digital twins and model context protocol." arXiv preprint arXiv:2506.13068 (2025).
17. Shilpashree, N., Girish Jadhav, and Abhijit Mitra. "Neuromorphic-Driven Agentic AI for Autonomous Decision-Making Systems." In 2024 4th International Conference on Mobile Networks and Wireless Communications (ICMNBC), pp. 1-8. IEEE, 2024.
18. Du, Changyu, Sebastian Esser, Stavros Nourias, and André Borrmann. "Text2BIM: Generating building models using a large language model-based multi-agent framework." arXiv preprint arXiv:2408.08054 (2024).
19. Jang, Suhyung, and Ghang Lee. "Interactive design by integrating a large pre-trained language model and building information modeling." In Computing in Civil Engineering 2023, pp. 291-299. 2023.
20. Samdani, Gaurav, Yawal Dixit, and Ganesh Viswanathan. "Agentic AI in autonomous financial advisories." World Journal of Advanced Engineering Technology and Sciences 9, no. 1 (2023): 410-420.
21. Sundaramurthy, Senthil Kumar, Nischal Ravichandran, Anil Chowdary Inaganti, and Rajendra Muppalaneni. "AI-powered operational resilience: Building secure, scalable, and intelligent enterprises." Artificial Intelligence and Machine Learning Review 3, no. 1 (2022): 1-10.
22. Gherhes, Cristian, Tim Vorley, Paul Vallance, and Chay Brooks. "The role of system-building agency in regional path creation: insights from the emergence of artificial intelligence in Montreal." Regional Studies 56, no. 4 (2022): 563-578.
23. Westman, Stina, Janne Kauttonen, Aarne Klemetti, Niilo Korhonen, Milja Manninen, Asko Mononen, Salla Niittymäki, and Henry Paananen. "Artificial Intelligence for Career Guidance--Current Requirements and Prospects for the Future." IAFOR Journal of Education 9, no. 4 (2021): 43-62.