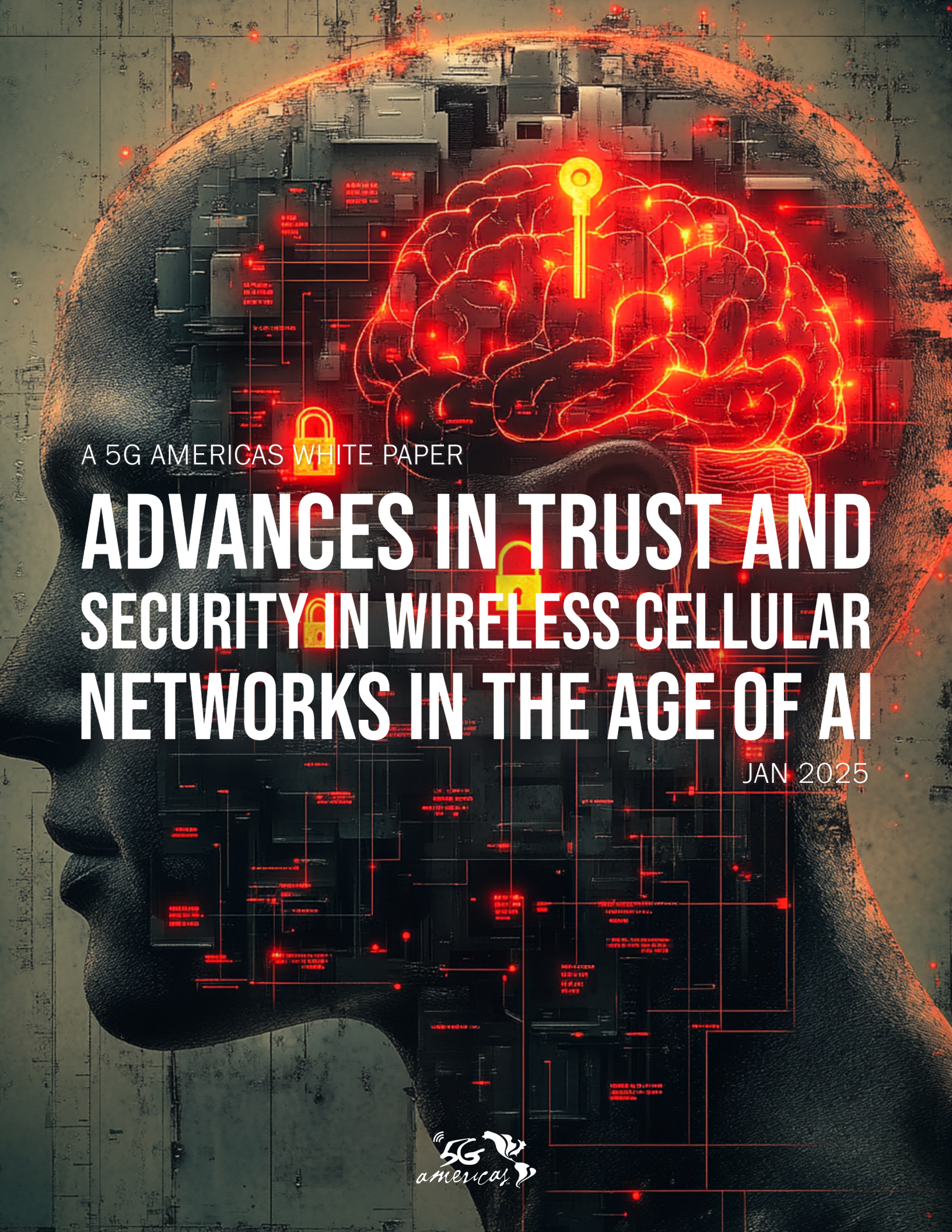




A 5G AMERICAS WHITE PAPER

ADVANCES IN TRUST AND SECURITY IN WIRELESS CELLULAR NETWORKS IN THE AGE OF AI

JAN 2025



Contents

Executive Summary..... 3

1. Artificial Intelligence and Machine Learning..... 4

 1.1 Introduction 4

 1.2 AI in Wireless Networks Today..... 4

 1.3 How AI Can be Leveraged in Wireless Networks..... 4

 1.4 Risk associated with AI 5

 1.5 Frameworks for Securing AI/ML..... 6

 1.6 AI-driven attacks on wireless networks: Impacts and controls 7

 1.7 Governance Around AI 8

 1.8 Recommendations 9

Conclusion 10

Appendix 11

Acknowledgments..... 12

Endnotes..... 13

Executive Summary

As technology continues to evolve, cellular wireless network organizations face the dual challenge of integrating new innovations while ensuring they remain trustworthy and secure. At the same time, the threat landscape is expanding, with malicious actors exploiting these advancements for harmful purposes. This paper provides a high-level overview of the upcoming security and trust challenges posed by Artificial Intelligence (AI).

Adoption of AI/ML based solutions has gained and continues to gain traction for diverse use cases and mobile networks are no exception. The majority of mobile network related AI/ML solutions to date can be considered proprietary, such as anomaly detection, performance improvements, and increased automation. However, studies to identify appropriate solutions to standardize aspects of AI/ML lifecycle management are currently ongoing. For example, 3GPP has already standardized one solution for AI/ML model training and inference in the 5G Core Network but ongoing studies continue. The O-RAN Alliance has also developed specifications related to AI/ML in the RIC (RAN Intelligent Controller) and has published a report on AI/ML security. Other international standardization bodies, such as ISO, have also been publishing standards to address risks specific to AI, both from organizational and product perspectives. For an in depth understanding of current and planned usage of AI in Cellular Networks readers are encouraged to read the 5G Americas white paper titled AI for Cellular Networks.¹

AI/ML solutions, including those used to protect networks, present an additional attack surface that an adversary can potentially target. On the other hand, an adversary could potentially use AI/ML as an attack vector to launch an attack on a network. It is therefore imperative that AI/ML assets used in mobile networks, such as training/test/validation data and trained models, and their associated parameters/hyperparameters, are protected from unauthorized access, tampering and theft. Equally important is that platforms, where AI/ML Assets are stored and/or processed, are secured in a robust manner. Currently available security best practices and frameworks can be leveraged with strong attention paid to securing AI/ML Assets and platforms.

Also increasing the threat landscape is the implementation of AI platforms. Although AI can potentially be used as a trust and security tool, AI must be designed, developed, and deployed in a secure way.

AI/ML platforms are used across wireless networks for a variety of tasking. Advances and investment in AI will increase its commonality in our wireless networks, functions, and tasking. However, the introduction of an AI platform comes at an increased risk. AI platforms are an attack vector and can also suffer from design and implementation failures. Utilization of secure by design practices and AI threat mapping and risk assessments are essential to ensure products and solutions are secure and risks are mitigated.

AI platforms, models, and/or data are often acquired from third party sources or vendors. Even when deploying these, it is essential organizations seek to thoroughly understand the risk and implement mitigations to ensure AI solutions are designed secure and maintained responsibility.

AI can also be used in conjunction with traditional wireless network attacks, such as eavesdropping, jamming, and spoofing. There are also new, and advanced attacks leverage AI on networks. We recommend investing not only in AI technologies, but in AI security solutions, to maintain pace with the changes threat landscape.

Regulations, frameworks, and standards are emerging globally to aid organizations in developing responsible and trustworthy AI. In the U.S., the Government and NIST also play an important role in aiding in the trust and security of AI platforms.

1. Artificial Intelligence and Machine Learning

1.1 Introduction

AI represents the simulation of human intelligence by machines that are programmed to think and learn like humans. These machines or systems can perform tasks that would typically require human intelligence. Some conventional examples include visual perception, speech recognition, decision-making, and language translation.

Machine Learning (ML) is a subset of AI that focuses on the development of algorithms and statistical models that enable computers to learn from and make predictions or decisions based on data. Instead of being explicitly programmed to perform a task, ML systems use patterns and inference to improve their performance over time. Common applications of ML include recommendation systems, fraud detection, and autonomous vehicles.

Generative AI (GenAI), a type of AI, enhances the capabilities of traditional ML by enabling machines to create new content, such as text, images, music, and even code. Unlike fixed purpose AI, which focuses on recognizing patterns and making decisions based on existing data, generative AI uses models to generate novel outputs. This capability makes it possible to automate creative tasks, generate realistic simulations, innovate, and streamline processes, and provide more personalized user experiences.

1.2 AI in Wireless Networks Today

AI/ML use cases in wireless networks today are limited, however, use cases for AI in telecommunication networks are evolving. AI/ML is actively being studied within 3GPP for the RAN, Core & O&M domains in different stages of completeness. Fundamental AI/ML framework aspects that are under study include where in the network data sets will be stored and which network entities are allowed access to these data sets, which network entities are authorized to perform training with which data sets and is training undertaken centrally, distributed or collaboratively, where trained models are stored, where inference is authorized to be undertaken and with which trained model and so on. For example, in the Core domain the NWDAF (Network Data Analytics Function) can assume the role of training and/or inference. Also, in the case where data sets must remain in a local domain for security and/or privacy reason and cannot be shared with a central domain for training

purposes, federated learning is supported whereby local NWDAF(s) trains the model locally and provides the locally trained model parameters to a central domain NWDAF which updated the central model. This enables the central model to be trained while still preserving the security/privacy of the local domain data. In the RAN domain an AI/ML framework and various use cases are under study, such as beam management and positioning optimization, as well as various model training scenarios between the UE and gNB such as one sided where the gNB only trains the model or collaborative where both the UE and gNB train the model in a coordinated manner.

1.3 How AI Can be Leveraged in Wireless Networks

AI can be leveraged in wireless networks in several new use cases. AI can help achieve efficiency gains, robust security measures, and advance network optimization when implemented safely and securely. Investments in the following use cases are expected:

AI Contact centers: By fine-tuning Large Language Models (LLMs) with telco-specific data, companies can create sophisticated digital assistants and chatbots. These AI tools can handle a wide range of customer interactions, from answering queries about subscriber plans to providing technical support for routers or set-top boxes. This approach allows for quick deployment of AI solutions, leveraging existing historical data to create personalized and efficient customer experiences.

Intelligent Network Planning and Self-Healing: By leveraging AI to process extensive network data, the system can predict needs, optimize configurations, and proactively address potential issues. This approach helps to streamline operations by reducing manual intervention, minimizing errors, and enabling predictive maintenance. It aims to ensure consistent network performance, prevent service disruptions, and improve overall network stability and reliability.

Power saving: By using machine learning techniques to predict data transmission patterns and network load, intelligent energy-saving strategies can be implemented. These strategies include adjusting Radio Access Network (RAN) configurations, such as cell activation/deactivation² and traffic offloading, dynamically configuring energy-saving parameters based on predicted load, and balancing system performance with energy efficiency.

Cell Load Balancing: AI/ML models are proposed to

address challenges in traditional load balancing methods, which often rely on current or past data and struggle with rapidly changing network conditions. By leveraging various measurements, historical data, and predictive analytics, AI-driven solutions aim to make more informed load balancing decisions.

Mobility Management: AI can be used to improve service continuity and user experience by minimizing call drops, radio link failures, unnecessary handovers, and ping-pong effects. AI/ML techniques are proposed to address challenges in conventional trial-and-error methods, particularly for applications with strict quality of service requirements.

AI for Security: Machine learning algorithms can analyze vast amounts of network traffic data in real-time to detect anomalies and potential cyber threats, allowing for rapid incident response. AI can be used to automate security patch management, predict vulnerabilities and prioritize updates to maintain robust defenses against evolving threats. In fraud detection, AI models can identify patterns indicative of SIM swapping, subscription fraud, or unauthorized account access.

AI Native 6G: Enhancements in the 6G physical layer design pursue better performance and energy efficiency, continuous improvement through lifetime machine learning model fine-tuning to adapt to local and changing conditions at both user equipment and network levels, and reduced research and development costs by using data-driven feature development instead of traditional algorithm design. These advancements enable efficient deployment of new AI-driven use cases, including those not yet envisioned, in particular, in intelligent Internet of Things, mixed reality, and collaborative autonomous agents. With significantly higher speeds, lower latency, and greater capacity, AI will be integral in optimizing radio networks of the future.

AI for vulnerability prediction, fraud & threat detection: Whether modeling a telecom network with Digital Twins or simulating system behavior from source code, AI/ML techniques can help strengthen security postures even more quickly. AI/ML algorithms can continuously monitor network traffic, examining data packets and user behavior to detect anomalous patterns or behaviors indicative of fraud or cyber threats. This permits early detection and prevention of activities such as unusual routing of calls, Subscriber Identity Module (SIM) card cloning, malware, denial of service attacks, and intrusions.

AI for Trust and Privacy: By implementing AI/ML responsibly, with robust accountability mechanisms like audits, disclosures, liability rules, and thoughtful guardrails, stakeholders can be confident that telecom AI systems will not harm them. While enabling valuable services, privacy-preserving ML techniques can also protect customer data.

1.4 Risk associated with AI

The integration of AI into the mobile network is expected to bring significant improvements in network performance and user experience. However, it also comes with potential risks, broadly categorised into Security, Regulatory, Trustworthiness and Sustainability. Some of these risks, for example, security, are common to all software systems, but exacerbated by the advance of AI. Other risks, like those related to trustworthiness, are directly related to the ever-improving quality of modern AI systems and are novel for humankind. Old or new, all the risks must be carefully considered and mitigated to ensure the benefits of ubiquitous use of AI outweigh the potential harms.

As AI systems become more prevalent in network management and operations, they also become attractive targets for cyberattacks. Security vulnerabilities in AI systems can arise at any stage of the lifecycle. For example, poisoning attacks occur at training time: attacker manipulates training data to introduce specific vulnerabilities to be exploited at deployment. In an adversarial attack, malicious actors manipulate input data at deployment to compromise network integrity and reliability. Evasion attacks are yet another threat: these use adversarial inputs or prompt injection to bypass security measures. In a black box attack, the attackers can steal, replicate, or reverse engineer a trained model (model extraction).

The energy consumption and computational resources required for training and deploying large AI models open yet another, novel, threat vector - a sponge attack. In this attack an adversary crafts model inputs to maximise model's energy consumption and latency, therefore driving user experience to the worst-case performance.³

Now that attackers are armed with more sophisticated AI, the pace of the security race has accelerated: the attacks evolve in real-time. While it is efficient to use AI in security processes to defend against cyberattacks, utmost care needs to be taken to make sure new attack vectors are not introduced.

While many AI applications aim to optimize energy efficiency, it is crucial to consider the rebound effect. As operations become more cost-effective, network traffic tends to increase, potentially negating any net decrease in energy consumption. Moreover, the energy demands associated with training and deploying large-scale AI models, compared to traditional fixed approaches, raise significant sustainability concerns.

The vast amounts of data captured by AI systems also increase the risk of data breaches and the need for increased security around data lakes. The use of sensitive data for AI training and inference, coupled with free-form user interface of new generative models are some of the factors that raise privacy concerns. Ensuring responsible data handling and compliance with privacy regulations throughout the AI system's lifecycle is paramount.

It is often difficult to explain and trust the decisions of AI systems due to their complexity ('the black box' challenge). 'Trustworthiness' encompasses several qualities related to trust, such as transparency, accountability, reliability, and fairness. Transparent AI systems provide all stakeholders with appropriate knowledge and skills to interpret the automated decisions. Accountability necessitates assigning responsibility for AI outcomes. Reliability and accuracy deal with correctness and consistent performance, which are necessary to create trust, although sometimes, overreliance on AI systems may occur, leading to further unexpected risks. Finally, fairness addresses concerns of biases inherent in training data, that may result in unfair treatment of certain user groups, thus eroding trust. These challenges are novel and difficult to address, and recently ethical AI frameworks and guidelines are being established globally to promote responsible AI development use. For example, ISO/IEC 42001:2023 *Artificial intelligence – Management System*,⁴ offers a framework to manage risks and opportunities associated with AI, from conception of the AI system to the end of service. Such standards can be used to balance caution with innovation while building trust with users and operators.

Navigating the risks associated with AI in the network requires a comprehensive approach that includes robust security measures, ethical considerations, and continuous monitoring and adaptation.

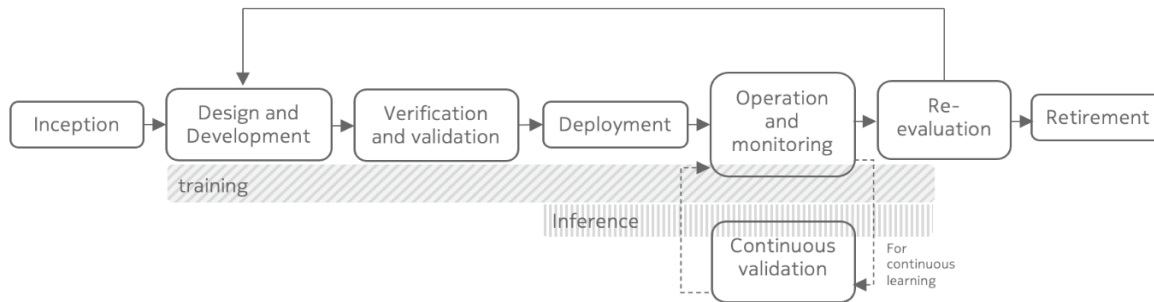
1.5 Frameworks for Securing AI/ML

Training and inference are the two key phases in Machine Learning Lifecycle. The training phase involves feeding data to a model and adjusting its parameters so that it can accurately make predictions on new, unseen data. Data used for the training phase is typically split into training data used to initially train the model and test/validation data used to subsequently verify and fine tune the model accuracy, with the training/testing/validation being an iterative process until the required level of accuracy is achieved. The inference phase involves using the trained model, with parameters learned during the training phase, to make predictions on new data it receives. A deployed system may be in both phases at once, as the system may learn continuously and re-train to adapt to shifts in real-world data.

In the context of lifecycle shown in Figure 2, it highlights key technical processes that occur during development of ML models:

- Definition of system requirements and model requirements
- Acquisition, update and preparation of data
- Training and re-training the model
- Verification of model before deployment
- Deployment of the model
- Continued testing of the model after deployment
- Retirement of the model

Figure 2. AI/ML lifecycle⁵



AI/ML adoption is growing rapidly, and it can positively affect every part of telecom systems. Systems that can create significant societal good could potentially be leveraged for harm. To prevent this, regulatory and technology leaders are looking to put in place guardrails to minimize bias, prevent unintended behavior, and maintain safe human interaction and consumption. AI/ML has tremendous potential for enhancing trust and security in telecom. Understanding the threat landscape to AI/ML is imperative before implementation. We recommend utilizing resources such as MITRE's *Adversarial Threat Landscape for AI Systems* (ATLAS)⁶. The Open Worldwide Application Security Project (OWASP) top ten list of security issues with ML⁷ and LLMs⁸ are also provide guidance on AI's latest known vulnerabilities.

AI/ML security can be applied to many areas, including open networks with increasingly large attack surfaces and complexity. We recommend utilizing a secure by design type of framework when designing, developing, and deploying AI. AI must be secured throughout its lifecycle. Securing AI lifecycles go beyond tradition Software Development Lifecycle, as AI has different threat vectors and resiliency needs. We also recommend utilizing a secure development framework that is AI/ML specific, such as MLSecOps.⁹

In 3GPP the SA3 Security group is evaluating the security and privacy threats and risks associated with AI/ML use and formulating consensus-based solution recommendations to address each issue. 3GPP has done a study in Release 18 and an ongoing study in Release 19 on AI/ML security. 3GPP SA3 concentrates on the interface security, and 3GPP interfaces are well secured via the existing security, i.e., mTLS, IPSec, etc. Regarding poisoning attacks and training attacks, these are considered implementation aspects and as such fall outside the scope of 3GPP but should be taken care of as part of secure software development & deployment lifecycle.

The O-RAN Alliance is also studying the security and privacy implications of AI/ML assets within the context of the O-RAN architecture and have published a technical report titled O-RAN Study on Security for AI/ML¹⁰ which identifies threats and risks and recommends potential security controls to protect against those threats through safeguards or mitigations.

In general, every stage of the AI/ML model training and deployment pipeline needs to be protected from unauthorized access, tampering & theft of the data sets, trained models and parameters.

1.6 AI-driven attacks on wireless networks: Impacts and controls

As discussed in the previous sections, AI is a powerful technology that can potentially enhance the performance, efficiency, and security of wireless networks. However, malicious actors could also use AI to launch sophisticated, intelligent, and adaptive cyberattacks on wireless networks. This includes 5G networks by using the same or similar AI techniques to learn from the network data and feedback, and to optimize the attack parameters and strategies. AI could also be used to generate realistic and adaptive fake signals, messages, or identities, by using generative adversarial networks, deep fakes, or synthetic data, to impersonate legitimate network users or devices, or create false network events or alarms. These attacks could have severe consequences on both the data that is handled, transferred, and stored by these components, and on the performance of the network itself

Cyberattacks on wireless networks are not a new concept being introduced by AI. However, AI-driven cyberattacks are different from the traditional attacks due to AI features and capabilities enabling threat actors to design and deploy complex and difficult to detect attack campaigns.

Additionally, AI could enable the attackers to coordinate and synchronize their actions, and to adapt and evolve their strategies by using multi-agent systems, game theory, or genetic algorithms, in response to the network security countermeasures.

Use of AI in conjunction with traditional wireless network attacks, such as eavesdropping, jamming, spoofing, etc. adds complexity to the execution and detection of attacks. In addition, AI has introduced new types of attacks such as intelligent jamming in 5G networks, enabling the attack tool to selectively and adaptively interfere with the 5G network communication by using machine learning or reinforcement learning techniques, based on the network data and feedback.

Similarly, on edge networks, attackers could use AI to create a new type of attack, called intelligent intrusion, to efficiently penetrate and compromise the edge network security system without detection. By using adversarial machine learning or reverse engineering, vulnerabilities or loopholes in the edge network security mechanisms, such as encryption, authentication, or firewall, and could be exploited in a targeted and efficient manner.

A non-exhaustive collection of features and capabilities of AI that are facilitating these changes in the threat landscape and supporting literature are listed below. However, it should be noted the degree which these threats are feasible in real-world deployments is not addressed.

- **The ability of AI to invoke attacks and adapt to configuration changes:** AI can autonomously launch cyberattacks on wireless networks by exploiting their vulnerabilities and weaknesses. AI can also adapt to the configuration changes and countermeasures by learning from their responses and modifying attack strategies accordingly.¹¹
- **Countering AI-driven attacks:** AI can make it difficult for the defenders to counter its real-time changes, analysis, and capabilities. AI can also use deception and obfuscation techniques to evade detection and attribution, by mimicking legitimate network traffic, hiding its malicious activities, or changing its identity and location.¹²
- **AI's ability to identify behavior patterns and vulnerabilities:** AI can use machine learning to analyze the behavior patterns and characteristics of wireless networks and their components, such as NG-RAN and 5GC. AI can then identify the soft targets and the optimal timing and methods to launch its attacks, by exploiting the network dynamics, congestion, latency, and load balancing.^{13 14}
- **Speed and destruction from AI attacks:** AI can potentially cause more damage and disruption to the wireless networks and data in a short window, by

using its speed, scalability, and parallelism. AI can also either hide or circumvent the traffic patterns, change the system log analysis process, or delete actionable data, to prevent the defenders from tracing and recovering from the attacks.^{15 16}

Due to these features and capabilities, AI-driven attacks could have severe impacts and consequences on the data that is handled, transferred, and stored by wireless network components and across the wireless networks in general. Any controls or solutions designed to prevent, detect, and mitigate AI-driven attacks on wireless networks could leverage the same features and capabilities to develop tools and solutions to counter AI-driven attacks and impacts.

1.7 Governance Around AI

The emergence and rapid rise in the use of AI and the increasing applications within telecommunications has created seemingly unlimited potential and at the same time, recognition of the potential risks. These risks highlight the need for clear AI governance across multiple layers including regulatory and legal approaches alongside industry and company level controls. Governance may span both national and international levels. Key stakeholders typically engaged are government agencies, industry bodies, telecommunications companies, AI developers and various advocacy groups. AI governance applied to the telecommunications sector requires a balanced approach focused on safeguarding privacy, security and ethical standards while still promoting innovation.

There are several key aspects of AI governance as it relates to telecommunications involving policies, regulations, and guidelines to ensure responsible and ethical use of AI. These span areas related to data privacy and governance, security, bias mitigation, transparency, accountability and compliance with existing laws and regulations.¹⁷

In the U.S., AI governance in telecommunications is a critical area of focus. There are several existing laws and regulations that apply to some aspects of AI governance that help to establish a foundation. These include:

- Federal Communications Commission's (FCC) role in areas of spectrum management, network security and consumer protection.
- Federal Trade Commission (FTC) Act provides the FTC with the authority to take action against unfair or deceptive practices in commerce, including those related to AI applications in telecommunications. Concerns like unfair bias or misleading representations in AI-driven telecommunications products or services may be addressed by this Act too.

- Privacy Regulations including the Communications Assistance for Law Enforcement Act (CALEA), Electronic Communications Privacy Act (ECPA), and regulations enforced by the FCC and FTC include provisions related to data privacy and security in telecommunications.
- Telecommunications Act of 1996 provides a regulatory framework for telecommunications services and may indirectly impact AI governance, while not addressing it directly.

Some newer developments in the U.S. that directly address AI governance include the following guidance, frameworks, and commitments. These efforts are examples of how the U.S. is actively working to establish robust AI governance for the telecom industry and beyond:

- White House Executive Order (EO): This EO includes measures to ensure safety, security and citizen rights in the content of AI, emphasizing safe, secure and trustworthy development and use of AI. Additionally, in July 2023, the White House secured voluntary commitments to safety, security, and responsible AI practices by the leading U.S. AI companies.¹⁸
- NIST AI Risk Management Framework 1.0: This framework, published in January 2023, provides guidelines for managing AI risks, ensuring transparency, and promoting accountability.¹⁹
- NIST has also established an AI Safety Institute to advance the science, practice, and adoption of AI safety across the spectrum of risks. The institute will focus first on the priorities of the White House EO noted above.
- Office of Management and Budget Guidance Memo: This memo outlines steps to enhance AI safety, security, and ethical use within the U.S. government.
- MIT Group's Policy Briefs: One of the main papers is "A Framework for U.S. AI Governance: Creating a Safe and Thriving AI Sector". Their goal being to enhance U.S. leadership in AI in general, while limiting harm that could result.²⁰

Telecommunication companies must establish a cohesive AI governance to mitigate the risks and challenges presented by AI in order to thrive in this dynamic industry. There are specific strategies that companies can utilize to mitigate the risks associated with AI use and address the challenges.

- Establish a comprehensive AI governance framework.
- Regularly review and update policies to align with advances in technology, industry standards and regulations.
- Ensure diverse stakeholder engagement, both within and outside of the organization.
- Implement tools and processes to identify and mitigate risks throughout the AI lifecycle. Such measures encompass every stage of development, including training data collection, development, testing and regular auditing of the AI models, deployment, and

post-market monitoring.

- Provide clarity and transparency to relevant stakeholders on AI use, sources, architecture, and decision-making processes.
- Facilitate AI literacy among stakeholders through training (for example, employees, users and wider public) to the extent needed to interpret the operation of the AI system and understand potential risks and mitigations, and to uphold ethical and safe practices.
- Collaborate with regulators and industry partners and peers to share best practices and address challenges.

A multi-stakeholder approach involving government bodies, regulators, industry organizations, telecommunication companies, and advocacy groups is crucial for effective AI governance. While regulators and government agencies develop policies to safeguard consumer privacy and protection, there's a risk that overly restrictive regulations could hinder innovation in the telecommunications sector. To mitigate this, we recommend the wireless industry proactively adopt robust internal AI governance frameworks that prioritize transparency, engagement, AI literacy, and collaboration to minimize potential risks and regulatory challenges.

1.8 Recommendations

AI/ML and its adoption into the telco domain is a rapidly evolving space as is the topic of securing AI/ML and therefore, the following recommendations are non-exhaustive with the potential to evolve. We recommend continuing to use and build upon baseline security controls (such as cryptographic controls, including PQC), supply chain security, secure software development, and the adoption of zero-trust principles for the protection of AI/ML assets (i.e. data sets, models, & parameters) from unauthorized access, tampering, and theft. Develop and deploy AI/ML in protected environments, including data storage and processing. Leverage AI/ML security frameworks, threat modeling, and secure lifecycle developments. Upskill and educate the workforce on AI/ML. For all stakeholders, proactively seek out and stay abreast of future AI telco security publications, recommendations, and new standards that may emerge.

Conclusion

This paper has discussed on a high-level the threat of AI/ML on wireless networks and outlined a number of recommendations to mitigate these respective threats. Understanding the risks and threats involved with implementing AI/ML is paramount in the securing of AI/ML in wireless networks. While AI/ML and the security of AI/ML is an evolving space, existing security best practices still provide a solid foundational layer to build upon complemented with AI/ML specific risk and governance frameworks. Governance plays a key role in trust and security, however, the responsibility up to all organizations to ensure the trust and security of wireless networks.

Appendix

Appendix A: Acronyms

AI: Artificial Intelligence	ML: Machine Learning
AS: Access Stratum	MNO: Mobile Network Operator
CA: Certificate Authority	NIST: National Institute of Standards and Technology
CALEA: Communications Assistance for Law Enforcement Act	OWASP: Open Worldwide Application Security Project
CFRG: Crypto Forum Research Group	PQC: Post-quantum cryptography
CISA: Cybersecurity and Infrastructure Security Agency	PQTN: Post-Quantum Telco Network
CRQC: Cryptographically Relevant Quantum Computer	PQUIP: Post-Quantum Use in Protocols
CSP: Communication Service Providers	QC: Quantum Cryptography
DTLS: Datagram Transport Layer Security	RAN: Radio Access Network
ECPA: Electronic Communications Privacy Act	RFC: Request for Comments
EO: Executive Order	SA: Security Aspects
FCC: Federal Communications Commission	SA: Security Associations
FTC: Federal Trade Commission	SBOM: Software Bill of Materials
GSMA: Global Mobile Supplier Association	SDO: Standards Development Organization
HSM: Hardware Security Modules	SIM: Subscriber Identity Module
IETF: Internet Engineering Task Force	SUCI: Subscription Concealed Identifier
IKE: Internet Key Exchange	SUPI: Subscriber Permanent Identity
IP: Internet Protocol	TLS: Transport Layer Security
LLM: Large Language Models	UE: User Equipment
MITRE: https://www.mitre.org/	ZTA: Zero-Trust Architecture

Acknowledgments

5G Americas facilitates and advocates for the advancement of 5G and beyond toward 6G throughout the Americas.

5G Americas' Board of Governors members include Airspan Networks, Antel, AT&T, Ciena, Cisco, Crown Castle, Ericsson, Liberty Latin America, Mavenir, Nokia, Qualcomm Incorporated, Rogers Communications, Samsung, T-Mobile USA, Inc., and Telefónica.

5G Americas would like to recognize the significant project leadership and important contributions of group leaders Taylor Hartley of Ericsson and Martin McGrath of Nokia - along with many representatives from member companies on 5G Americas' Board of Governors who participated in the development of this white paper.

The contents of this document reflect the research, analysis, and conclusions of 5G Americas and may not necessarily represent the comprehensive opinions and individual viewpoints of each particular 5G Americas member company. 5G Americas provides this document and the information contained herein for informational purposes only, for use at your sole risk. 5G Americas assumes no responsibility for errors or omissions in this document. This document is subject to revision or removal at any time without notice. No representations or warranties (whether expressed or implied) are made by 5G Americas and 5G Americas is not liable for and hereby refuses any direct, indirect, punitive, special, incidental, consequential, or exemplary damages arising out of or in connection with the use of this document and any information contained in this document.

© Copyright 2024 5G Americas

Endnotes

- 1 5G Americas, Artificial Intelligence in Cellular Networks
- 2 <https://ieeexplore.ieee.org/abstract/document/10188338>
- 3 [Sponge Examples: Energy-Latency Attacks on Neural Networks | IEEE Conference Publication | IEEE Xplore](#)
- 4 [ISO/IEC 42001:2023 - AI management systems](#)
- 5 [ISO/IEC 22989:2022 - Information technology — Artificial intelligence — Artificial intelligence concepts and terminology](#)
- 6 <https://atlas.mitre.org/>
- 7 <https://owasp.org/www-project-machine-learning-security-top-10/>
- 8 <https://owasp.org/www-project-top-10-for-large-language-model-applications/>
- 9 [MLSecOps | Home](#)
- 10 <https://specifications.o-ran.org/download?id=699>
- 11 [Buczak & Guven, 2016: A Survey of Data Mining and Machine Learning Methods for Cyber Security Intrusion Detection | IEEE Journals & Magazine | IEEE Xplore](#)
- 12 [Wang et al., 2019: With Great Training Comes Great Vulnerability: Practical Attacks against Transfer Learning | USENIX](#)
- 13 [Alshamrani et al., 2019: dblp: A Survey on Advanced Persistent Threats: Techniques, Solutions, Challenges, and Research Opportunities.](#)
- 14 [Zhang et al., 2020: Artificial intelligence in recommender systems | Complex & Intelligent Systems \(springer.com\)](#)
- 15 [Bajpai et al., 2019: Artificial Intelligence, the Law and the Future by G. S. Bajpai :: SSRN](#)
- 16 [Li et al., 2019: AI | Free Full-Text | Ethics & AI: A Systematic Review on Ethical Concerns and Related Strategies for Designing with AI in Healthcare \(mdpi.com\)](#)
- 17 [The Oxford Handbook of AI Governance | Oxford Academic \(oup.com\)](#)
- 18 [Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence | The White House](#)
- 19 <https://doi.org/10.6028/NIST.AI.600-1>
- 20 [AIPolicyBrief.pdf \(mit.edu\)](#)